

VISIT...

LANZAROTE
Caliente.COM

AETHER MAGE



ERIC DOLLARD COLLECTION

ÆTHERFORCE

GENERAL LECTURES

ON

ELECTRICAL ENGINEERING

BY

CHARLES PROTEUS STEINMETZ, A. M., Ph. D.

Consulting Engineer of the General Electric Company,
Professor of Electrical Engineering in Union University,
Past President, A. I. E. E.

Author of

"Alternating Current Phenomena,"
"Elements of Electrical Engineering,"
"Transient Electric Phenomena and Oscillations."

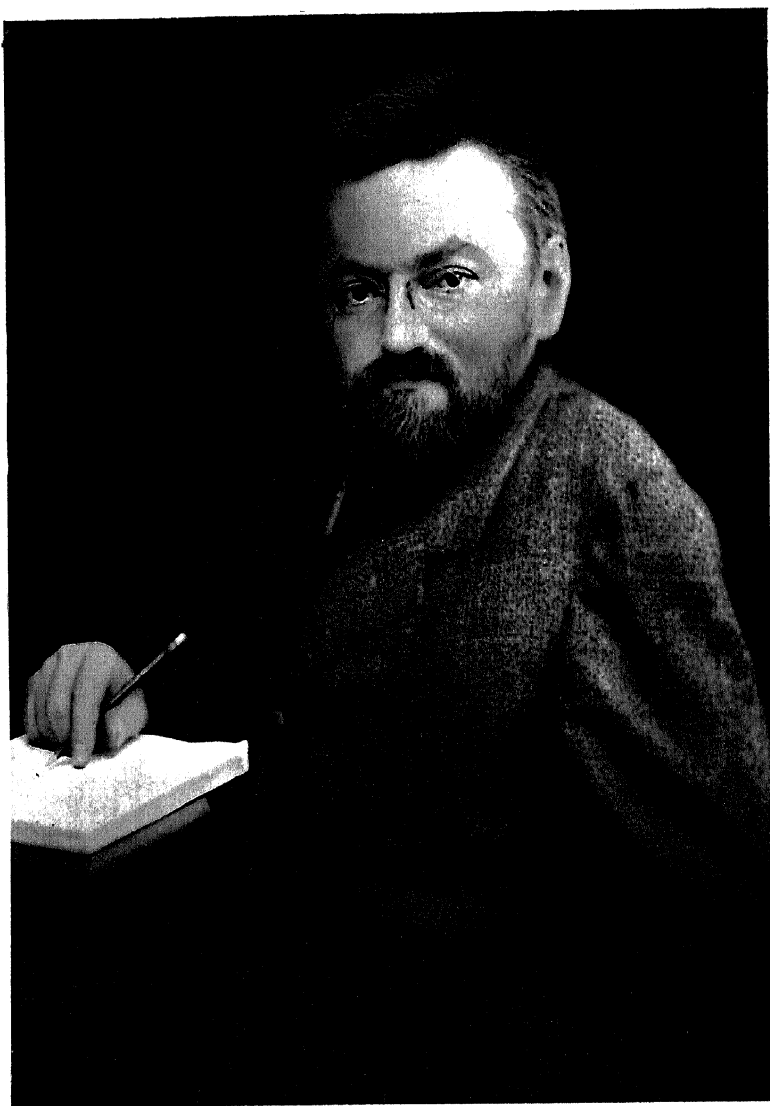
Second Edition.

Compiled and Edited by

JOSEPH Le ROY HAYDEN

Robson & Adey, Publishers
Schenectady, N. Y.

ÆTHERFORCE



ÆTHERFORCE

Copyright 1908 by
Robson & Adee

Contents

First Lecture—General Review.....	7
Second Lecture—General Distribution.....	21
Third Lecture—Light and Power Distribution.....	35
Fourth Lecture—Load Factor and Cost of Power....	49
Fifth Lecture—Long Distance Transmission.....	61
Sixth Lecture—Higher Harmonics of the Generator Wave.....	77
Seventh Lecture—High Frequency Oscillations and Surges.....	89
Eighth Lecture—Generation.....	99
Ninth Lecture—Hunting of Synchronous Machines..	113
Tenth Lecture—Regulation and Control.....	125
Eleventh Lecture—Lightning Protection.....	135
Twelfth Lecture—Electric Railway.....	147
Thirteenth Lecture—Electric Railway Motor Char- acteristics..	163
Fourteenth Lecture—Alternating Current Railway Motors.....	175
Fifteenth Lecture—Electrochemistry.....	197
Sixteenth Lecture—The Incandescent Lamp.....	207
Seventeenth Lecture—Arc Lighting.....	215
Appendix I.—Light and Illumination.....	229
Appendix II.—Lightning and Lightning Protection..	259

Preface

THE following lectures on Electrical Engineering are general in their nature, dealing with the problems of generation, control, transmission, distribution and utilization of electric energy; that is, with the operation of electric systems and apparatus under normal and abnormal conditions, and with the design of such systems; but the design of apparatus is discussed only so far as it is necessary to understand their operation, and so judge of their proper field of application.

Due to the nature of the subject, and the limitations of time and space, the treatment had to be essentially descriptive, and not mathematical. That is, it comprises a discussion of the different methods of application of electric energy, the means and apparatus available, the different methods of carrying out the purpose, and the relative advantages and disadvantages of the different methods and apparatus, which determine their choice.

It must be realized, however, that such a discussion can be general only, and that there are, and always will be, cases in which, in meeting special conditions,, conclusions regarding systems and apparatus may be reached, differing from those which good judgment would dictate under general and average conditions. Thus, for instance, while certain transformer connections are unsafe and should in general be avoided, in special cases it may be found that the danger incidental to their use is so remote as to be overbalanced by some advantages which they may offer in the special case, and their use would thus be

PREFACE

justified in this case. That is, in the application of general conclusions to special cases, judgment must be exerted to determine, whether, and how far, they may have to be modified. Some such considerations are indicated in the lectures, others must be left to the judgment of the engineer.

The lectures have been collected and carefully edited by my assistant, Mr. J. L. R. Hayden, and great thanks are due to the publishers, Messrs. Robson & Adey, for the very creditable and satisfactory form in which they have produced the book.

CHARLES P. STEINMETZ.

Schenectady, N. Y., Sept. 5, 1908.

FIRST LECTURE

GENERAL REVIEW

IN ITS economical application, electric power passes through the successive steps : generation, transmission, conversion, distribution and utilization. The requirements regarding the character of the electric power imposed by the successive steps, are generally different, frequently contradictory, and the design of an electric system is therefore a compromise. For instance, electric power can for most purposes be used only at low voltage, 110 to 600 volts, while economical transmission requires the use of as high voltage as possible. For many purposes, as electrolytic work, direct current is necessary; for others, as railroading, preferable; while for transmission, alternating current is preferable, due to the great difficulty of generating and converting high voltage direct current. In the design of any of the steps through which electric power passes, the requirements of all the other steps so must be taken into consideration. Of the greatest importance in this respect is the use to which electric power is put, since it is the ultimate purpose for which it is generated and transmitted; next in importance is the transmission, as the long distance transmission line usually is the most expensive part of the system, and in the transmission the limitation is more severe than in any other step through which the electric power passes.

The main uses of electric power are :

General Distribution for Lighting and Power. The relative proportion between power use and lighting may vary from the distribution system of many small cities, in which

practically all the current is used for lighting, to a power distribution for mills and factories, with only a moderate lighting load in the evening.

The electric railway.

Electrochemistry.

For convenience, the subject will be discussed under the subdivisions:

1. General distribution for lighting and power.
2. Long distance transmission.
3. Generation.
4. Control and protection.
5. Electric railway.
6. Electrochemistry.
7. Lighting.

CHARACTER OF ELECTRIC POWER.

Electric power is used as—

- a. Alternating current and direct current.
- b. Constant potential and constant current.
- c. High voltage and low voltage.

a. Alternating current is used for transmission, and for general distribution with the exception of the centers of large cities; direct current is usually applied for railroading. For power distribution, both forms of current are used; in electrochemistry, direct current must be used for electrolytic work, while for electric furnace work alternating current is preferable.

The two standard frequencies of alternating current are 60 cycles and 25 cycles. The former is used for general distribution for lighting and power, the latter for conversion to direct current, for alternating current railways, and for large powers.

In England and on the continent, 50 cycles is standard frequency. This frequency still survives in this country in Southern California, where it was introduced before 60 cycles was standard.

The frequencies of 125 to 140 cycles, which were standard in the very early days, 20 years ago, have disappeared.

The frequency of 40 cycles, which once was introduced as compromise between 60 and 25 cycles is rapidly disappearing, as it is somewhat low for general distribution, and higher than desirable for conversion to direct current. It was largely used also for power distribution in mills and factories as the lowest frequency at which arc and incandescent lighting is still feasible; for the reason that 40 cycle generators driven by slow speed reciprocating engines are more easily operated in parallel, due to the lower number of poles. With the development of the steam turbine as high speed prime mover, the conditions in this respect have been reversed, and 60 cycles is more convenient, giving more poles at the same generator speed, and so less power per pole.

Sundry odd frequencies, as 30 cycles, 33 cycles, 66 cycles, which were attempted at some points, especially in the early days, have not spread; and frequencies below 25 cycles, as 15 cycles and 8 cycles, as proposed for railroading, have not proved of sufficient advantage—at least not yet—so that in general, in the design of an electric system, only the two standard frequencies, 25 and 60 cycles, come into consideration.

b. Constant current, either alternating or direct, that is, a current of constant amperage, varying in voltage with the load, is mostly used for street lighting by arc lamps; for all other purposes, constant potential is employed.

c. For long distance transmission, the highest permissible voltage is used; for primary distribution by alternating current, 2200 volts, that is, voltages between 2000 and 2600; for alternating current secondary distribution, and direct current distribution, 220 to 260 volts, and for direct current railroading, 550 to 600 volts.

I. GENERAL DISTRIBUTION FOR LIGHTING AND POWER.

In general distribution for lighting and power, direct current and 60 cycles alternating current are available. 25 cycles alternating current is not well suited, since it does not permit arc lighting, and for incandescent lighting it is just at the limit, where under some conditions and with some generator waves, flickering shows, while with others it does not show appreciably.

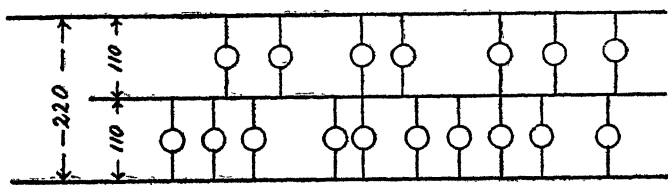


Fig. 1

The distribution voltage is determined by the limitation of the incandescent lamp, as from 104 to 130 volts, or about 110 volts. 110 volts is too low to distribute with good regulation, that is, with negligible voltage drop, any appreciable amount of power, and so practically always twice that voltage is employed in the distribution, by using a three-wire system, with 110 volts between outside and neutral, and 220 volts between the outside conductors, as shown diagrammatically in Fig. 1. By approximately balancing the load between the two circuits, the current in the neutral conductor is very small, the

drop of voltage so negligible, and the distribution, regarding voltage drop and copper economy, so takes place at 220 volts, while the lamps operate at 110 volts. Even where a separate transformer feeds a single house, usually a three-wire distribution is preferable, if the number of lamps is not very small.

When speaking of a distribution voltage of 110, some voltage anywhere in the range from 104 to 130 volts is employed. Exactly 110 volts is rarely used, but the voltages of distribution systems in this country are distributed over the whole range, so as to secure best economy of the incandescent lamp.

This condition was brought about by the close co-operation, in this country, between the illuminating companies and the manufacturers of incandescent lamps. The constants of an incandescent lamp are the candle power—for instance 16; the economy—for instance 3.1 watts for horizontal candle power; and the voltage—for instance 110. By careful manufacture, a lamp can be made in which the filament reaches 3.1 watts per candle power economy at 16 c. p. within one-half candle-power; but the attempt to fulfill at the same time the condition, that this economy and candle power be reached at 110 volts, within one-half volt, would lead to a considerable percentage of lamps which would fall outside of the narrow range permitted in the deviation from the three constants; and so, if the same distribution voltage were used throughout the country, either a much larger margin of variation would have to be allowed in the product, that is, the lamps would be far less uniform in quality—as is the case abroad,—or a large number of lamps would not fulfill the requirements, could not be used, and so would increase the cost of the rest.

Therefore, all the efforts in manufacture are concentrated on producing the specified candle power at the required economy, and the lamps are then sorted for voltage. This arrangement scatters the lamps over a considerable voltage range, and different voltages are then adopted by different distribution systems, so as to utilize the entire product of manufacture at its maximum economy. The result of this co-operation between lamp manufacturers and users is, that the incandescent lamps are very much closer to requirements, and more uniform, than would be possible otherwise. The effect however is, that the distribution is rarely actually 110, and in alternating current systems, the primary distribution voltage not 2200, but some voltage in the range between 2080 and 2600, as in step-down transformers a constant ratio of transformation, of a multiple of $10 \div 1$, is always used.

In the following, therefore, when speaking of 110, 220 or 2200 volts in distribution systems, always one of the voltages within the range of the lamp voltages is understood.

In this country, 110 volt lamps are used almost exclusively, while in England, for instance, 220 volt lamps are generally used, in a three-wire distribution system with 440 volts between the outside conductors. The amount of copper required in the distribution system, with the same loss of power in the distributing conductors, is inversely proportional to the square of the voltage. That is, at twice the voltage, twice the voltage drop can be allowed for the same distribution efficiency; and as at double voltage the current is one-half, for the same load twice the voltage drop at half the current gives four times the resistance, that is, one-quarter the conductor material. By the change from the 220 volt distribution with 110 volt lamps, to the 440 volt distribution with 220 volt

lamps, the amount of copper in the distributing conductor, and thereby the cost of investment can be greatly reduced, and current supplied over greater distances, so that from the point of view of the economical supply of current at the customers' terminals, the higher voltage is preferable. However, in the usual sizes, from 50 to 60 watts power consumption and so 16 candle power with the carbon filament, and correspondingly higher candle power with the more efficient metallized carbon and metal filaments, the 220 volt lamp is from 10 to 15% less efficient, that is, requires from 10 to 15% more power than the 110 volt lamp, when producing the same amount of light at the same useful life. This difference is inherent in the incandescent lamp, and is due to the far greater length and smaller section of the 220 volt filament, compared with the 110 volt filament, and therefore no possibility of overcoming it exists; if it should be possible to build a 220 volt 16 candle power lamp as efficient—at the same useful life of 500 hours—as the present 110 volt lamp, this would simply mean, that by the same improvement the efficiency of the 110 volt lamp could also be increased from 10 to 15%, and the difference would remain. For smaller units than 16 candle power, the difference in efficiency is still greater.

This loss of efficiency of 10 to 15%, resulting from the use of the 220 volt lamp, is far greater than the saving in power and in cost of investment in the supply mains; and the 220 volt system with 110 volt lamps is therefore more efficient, in the amount of light produced in the customer's lamps, than the 440 volt system with 220 volt lamps. In this country, since the early days, the illuminating companies have accepted the responsibility up to the output in light at the customer's lamps, by supplying and renewing the lamps free of charge, and the system using 110 volt lamps is therefore universally

employed while the 220 volt lamp has no right to existence; while abroad, where the supply company considers its responsibility ended at the customer's meter, and the customer is left to supply his own lamps, the supply company saves by the use of 440 volt systems—at the expense of a waste of power in the customer's 220 volt lamps, far more than the saving effected by the supply company.

In considering distribution systems, it therefore is unnecessary to consider any other lamp voltage than 110 volts (that is, the range of voltage represented thereby).

In direct current distribution systems, as used in most large cities, the 220 volt network is fed from a direct current generating station, or—as now more frequently is the case—from a converter substation, which receives its power as three-phase alternating, usually 25 cycles, from the main generating station, or long distance transmission line. In alternating current distribution, the 220 volt distribution circuits are fed by step-down transformers from the 2200 volt primary distribution system. In the latter case, where considerable motor load has to be considered, some arrangement of polyphase supply is desirable, as the single-phase motor is inferior to the polyphase motor, and so the latter is preferable for large and moderate sizes.

COMPARISON OF ALTERNATING CURRENT AND DIRECT CURRENT

At the low distribution voltage of 220, current can economically be supplied from a moderate distance only, rarely exceeding from 1 to 2 miles. In a direct current system, the current must be supplied from a generating station or a converter substation, that is, a station containing revolving machinery. As such a station requires continuous atten-

tion, its operation would hardly be economical if not of a capacity of at least some hundred kilowatts. The direct current distribution system therefore can be used economically only if a sufficient demand exists, within a radius of 1 to 2 miles, to load a good sized generator or converter substation. The use of direct current is therefore restricted to those places where a fairly concentrated load exists, as in large cities; while in the suburbs, and in small cities and villages, where the load is too scattered to reach from one low tension supply point, sufficient customers to load a substation, the alternating current must be used, as it requires merely a step-down transformer which needs no attention.

In the interior of large cities, the alternating current system is at a disadvantage, because in addition to the voltage consumed by resistance, an additional drop of voltage occurs by self-induction, or by reactance; and with the large conductors required for the distribution of a large low tension current, the drop of voltage by self-induction is far greater than that by resistance, and the regulation of the system therefore is seriously impaired, or at least the voltage regulation becomes far more difficult than with direct current. A second disadvantage of the alternating current for distribution in large cities is, that a considerable part of the motor load is elevator motors, and the alternating current elevator motor is inferior to the direct current motor. Elevator service essentially consists in starting at heavy torque, and rapid acceleration, and in both of these features the direct current motor with compound field winding is superior, and easier to control.

Where therefore direct current can be used in low tension distribution, it is preferable to use it, and to relegate alternating current low tension distribution to those cases where direct

current cannot be used, that is, where the load is not sufficiently concentrated to economically operate converter substations.

The loss of power in the low tension direct current system is merely the i^2r loss in the conductors, which is zero at no load, and increases with the load; the only constant loss in a direct current distribution system is the loss of power in the potential coils of the integrating wattmeters on the customer's premises. In the direct current system therefore, the efficiency of distribution is highest at light load, and decreases with increasing load.

In an alternating current distribution system, with a 2200 volt primary distribution, feeding secondary low tension circuits by step-down transformers, the i^2r loss in the conductors usually is far smaller than in the direct current system, but a considerable constant, or "no load", loss exists; the core-loss in the transformers, and the efficiency of an alternating current distribution is usually lowest at light load, but increases with increase of load, since with increasing load the transformer core loss becomes a lesser and lesser percentage of the total power. The i^2r loss in alternating current systems must be far lower than in direct current systems:

1. Because it is not the only loss, and the existence of the "no load" or transformer core loss requires to reduce the load loss or i^2r loss, if an equally good efficiency is desired. With an alternating current system, each low tension main requires only a step-down transformer, which needs no attention; therefore many more transformers can be used than rotary converter substations in a direct current system, and the i^2r loss is then reduced by the greatly reduced distance of secondary distribution.

2. In the alternating current system, the drop of voltage in the conductors is greater by the self-inductive drop than the

ir drop; the ir drop is therefore only a part of the total voltage drop; and with the same voltage drop and therefore the same regulation as a direct current system, the i^2r loss in the alternating current system would be smaller than in the direct current system.

3. Due to the self-inductive drop, smaller and therefore more numerous low tension distribution circuits must be used with alternating current than with direct current, and a separate and independent voltage regulation of each low tension circuit—that is—each transformer, therefore usually becomes impracticable. This means that the total voltage drop, resistance and inductance, in the alternating current low tension distribution circuits must be kept within a few percent, that is, within the range permissible by the incandescent lamp. As a result thereof, the voltage regulation of an alternating current low tension distribution is usually inferior to that of the direct current distribution—in many cases to such an extent as to require the use of incandescent lamps of lower efficiency. While therefore in direct current distribution 3.1 watt lamps are always used, in many alternating current systems 3.5 watt lamps have to be used, as the voltage regulation is not sufficiently good to get a satisfactory life from the 3.1 watt lamps.

SECOND LECTURE

GENERAL DISTRIBUTION

DIRECT CURRENT DISTRIBUTION

THE TYPICAL direct current distribution is the system of feeders and mains, as devised by Edison, and since used in all direct current distributions. It is shown diagrammatically in Fig. 2. The conductors are usually under-

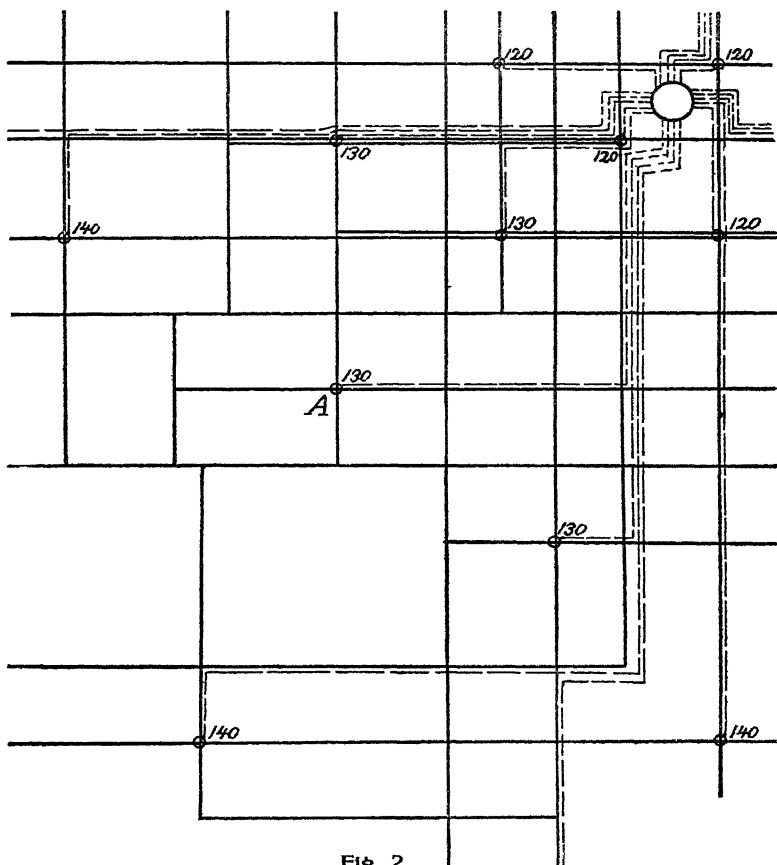


Fig. 2

ground, as direct current systems are used only in large cities. A system of three-wire conductors, called the "mains" is laid in the streets of the city, shown diagrammatically by the heavily drawn lines. Commonly, conductors of one million circular mil section (that is, a copper section which as solid round conductor would have a diameter of 1") are used for the outside conductors, the "positive" and the "negative" conductor; and a conductor of half this size for the middle or "neutral" conductor. The latter is usually grounded, as protection against fire risk, etc. Conductors of more than one million circular mils are not used, but when the load exceeds the capacity of such conductors, a second main is laid in the same street. A number of feeders, shown by dotted lines in Fig. 2, radiate from the generating station or converter substations, and tap into the mains at numerous points; potential wires run back from the mains to the stations, and so allow of measuring, in the station, the voltage at the different points of the distribution system. All the customers are connected to the mains, but none to the feeders. The mains and feeders are arranged so that no appreciable voltage drop takes place in the mains, but all drop of voltage occurs in the feeders; and as no customers connect to the feeders, the only limit to the voltage drop in the feeders is efficiency of distribution. The voltage at the feeding points into the mains is kept constant by varying the voltage supply to the feeders with the changes of the load on the mains. This is done by having a number of outside bus bars in the station, as shown diagrammatically in Fig. 3, differing from each other in voltage, and connecting feeders over from bus bar to bus bar, with the change of load.

For instance, in a 2 x 120 voltage distribution, the station may have, in addition to the neutral bus bar zero, three positive

bus bars 1, 1', 1'', and three negative bus bars 2, 2', 2'', differing respectively from the neutral bus by 120, 130 and 140 volts, as shown in Fig. 3. At light load, when the drop of voltage in the feeders is negligible, the feeders connect to the busses 1, 0, 2 of 120 volts. When the load increases, some of the feeders are shifted over, by transfer bus bars, to the 130 volt busbars 1' and 2'; with still further increase of load, more feeders are connected over to 130 volts; then some feeders are connected to the 140 volt bus bars, 1'' and 2'', and so, by varying

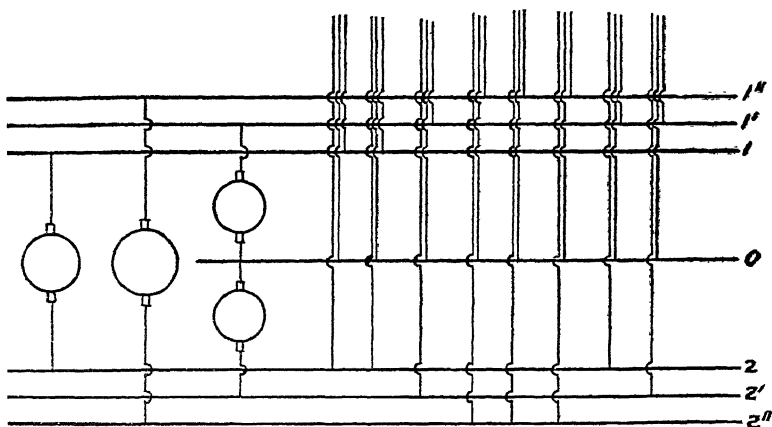


Fig. 3

the voltage supply to the feeders, the voltage at the mains can be maintained constant with an accuracy depending on the number of bus bars. It is obvious that a shift of a feeder from one voltage to another does not mean a corresponding voltage change on the main supplied by it, but rather a shift of load between the feeders, and so a readjustment of the total voltage in the territory near the supply point of the feeder. For instance, if by the potential wires a drop of voltage below 120 volts is registered in the main at the connection point of feeder A in Fig. 2, and this feeder then shifted from the supply

voltage 130 to 140, the current in the main near A, which before flowed towards A as minimum voltage point, reverses in direction, flows away from A, the load on feeder A and therefore increases, and the drop of voltage in A increases, while the load on the adjacent feeders decreases, and thereby their drop of voltage decreases, with the result of bringing up the voltage in the mains at the feeder A and all adjacent feeders. This inter-linkage of feeders therefore allows a regulation of voltage in the mains, far closer than the number of voltages available in the station.

The different bus bars in the station are supplied with their voltage by having different generators or converters in the station operate at different voltages, and with increasing load on the station, and consequent increasing demand of higher voltage by the feeders, shift machines from lower to higher voltage bus bars, inversely with decreasing load; or the different bus bars are operated through boosters, or by connection with the storage battery reserve, etc.

In addition to feeders and mains, tie feeders usually connect the generating station or substation with adjacent stations, so that during periods of light load, or in case of breakdown, a station may be shut down altogether and supplied from adjacent stations by tie feeders. Such tie feeders also permit most stations to operate without storage battery reserve, that is, to concentrate the storage batteries in a few stations, from which in case of a breakdown of the system, the other stations are supplied over the tie feeders.

ALTERNATING CURRENT DISTRIBUTION

The system of feeders and mains allows the most perfect voltage regulation in the distributing mains. It is however applicable only to direct current distribution in a territory of

very concentrated load, as in the interior of a large city, since the independent voltage regulation of each one of numerous feeders is economically permissible only where each feeder represents a large amount of power; with alternating current systems, the inductive drop forbids the concentration of such large currents in a single conductor. That is, conductors of one million circular mils cannot be used economically in an alternating current system.

The resistance of a conductor is inversely proportional to the size or section of the conductor, hence decreases rapidly with increasing current: a conductor of one million circular mils is one-tenth the resistance of a conductor of 100,000 circular mils, and so can carry ten times the direct current with the same voltage drop. The reactance of a conductor, however, and so the voltage consumed by self-induction, decreases only very little with the increasing size of a conductor, as seen from the table of resistances and reactances of conductors. A wire No. 000 B & S G is eight times the section of a wire No. 7, and therefore one-eighth the resistance; but the wire No. 000 has a reactance of .109 ohms per 1000 feet, the wire No. 7 has a reactance of .133 oms, or only 1.22 times as large. Hence, while in the wire No. 7, the reactance, at 60 cycles, is only .266 times the resistance and therefore not of serious importance, in a wire No. 000 the reactance is 1.76 times the resistance, and the latter conductor is likely to give a voltage drop far in excess of the ohmic resistance drop. The ratio of reactance to resistance therefore rapidly increases with increasing size of conductor, and for alternating currents, large conductors cannot therefore be used economically where close voltage regulation is required.

With alternating currents it therefore is preferable to use several smaller conductors in multiple: two conductors of

No. 1 in multiple have the same resistance as one conductor of No. 000; but the reactance of one conductor No. 000 is .109 ohms, and so 1.88 times as great as the reactance of two conductors of No. 1 in multiple, which latter is half that of one conductor No. 1, or .058 ohms, provided that the two conductors are used as separate circuits.

In alternating current low tension distribution, the size of the conductor and so the current per conductor, is limited by the self-inductive drop, and alternating current low tension networks are therefore of necessity of smaller size than those of direct current distribution.

As regards economy of distribution, this is not a serious objection, as the alternating current transformer and primary distribution permits the use of numerous secondary circuits.

In alternating current systems, a primary distribution system of 2200 volts is used, feeding step-down transformers.

The different arrangements are—

a. A separate transformer for each customer. This is necessary in those cases where the customers are so far apart from each other that they cannot be reached by the same low tension or secondary circuit; every alternating current system therefore has at least a number of instances where individual transformers are used.

This is the most uneconomical arrangement. It requires the use of small transformers, which are necessarily less efficient and more expensive per kilowatt, than large transformers. The transformer must be built to carry, within its overload capacity, all the lamps installed by the customer, since all the lamps may be used occasionally. Usually, however, only a small part of the lamps are in use, and those only for a small part of the day; so that the average load on the transformer is a very small part of its capacity.

As the core loss in the transformer continues whether the transformer is loaded or not, but is not paid for by the customer, the economy of the arrangement is very low; and so it can be understood that in the early days, where this arrangement was generally used, the financial results of most alternating current distributions were very discouraging.

Assuming as an instance a connected load of twenty 16 candle power lamps—low efficiency lamps, of 60 watts per lamp, since the voltage regulation cannot be very perfect—allowing then in cases of all lamps being used, an overload of 100%, which is rather beyond safe limits, and permissible only on the assumption that this load will occur very rarely, and for a short time—the transformer would have 600 watt rating. Assuming a core loss of 4%, this gives a continuous power consumption of 24 watts. Usually probably only one or two lamps will be burning, and these only a few hours per day, so that the use of two lamps, at an average—summer and winter—of three hours per day, would probably be a fair example of many such cases. Two lamps or 120 watts, for three hours per day, give an average power of 15 watts, which is paid for by the customer, while the continuous loss in the transformer is 24 watts; so that the all year efficiency, or the ratio of the power paid for by the customer, to the power consumed by the transformer, is only $\frac{15}{15 + 24}$ or 38%.

By connecting several adjacent customers to the same transformer, the conditions immediately become far more favorable. It is extremely improbable that all the customers will burn all their lamps at the same time, the more so, the greater the number of customers is, which are supplied from the same transformer. It therefore becomes unnecessary to

allow a transformer capacity capable of operating all the connected load. The larger transformer also has a higher efficiency. Assuming therefore as an instance, four customers of 20 lamps connected load each. The average load would be about 8 lamps. Assuming even one customer burning all 20 lamps, it is not probable that the other customers together would at this time burn more than 10 to 15 lamps, and a transformer carrying 30 to 35 lamps at overload would probably be sufficient. A 1500 watt transformer would therefore be larger than necessary. At 3% coreloss, this gives a constant loss of 45 watts, while an average load of 8 lamps for 3 hours per day gives a useful output of 60 watts, or an all year efficiency of nearly 60%, while a 1000 watt transformer would give an all year efficiency of 67%.

This also illustrates that in smaller transformers a low coreloss is of utmost importance, while the i^2r loss is of very secondary importance, since it is appreciable only at heavy load, and therefore affects the all year efficiency very little.

When it becomes possible to connect a large number of customers to a secondary main fed from one large transformer the connected load ceases to be of moment in the transformer capacity; the transformer capacity is determined by the average load, with a safe margin for overloads; in this case, good all year efficiencies can be reached.

Economical alternating current distribution therefore requires the use of secondary distribution mains of as large an extent as possible, fed by large transformers. The distance, however, to which a transformer can supply secondary current, is rather limited by the inductive drop of voltage; therefore, for supplying secondary mains, transformers of larger size than 30 kw. are rarely used, but rather several transformers are employed, to feed in the same main at different points.

Extending the secondary mains still further by the use of several transformers feeding into the same mains, or, as it may be considered, inter-connecting the secondary mains of the different transformers, we arrive at a system somewhat similar to the direct current system: a low tension distribution system of 220 volts three-wire mains, with a system of feeders tapping into it at a number of points, as shown in Fig. 4. These feeders

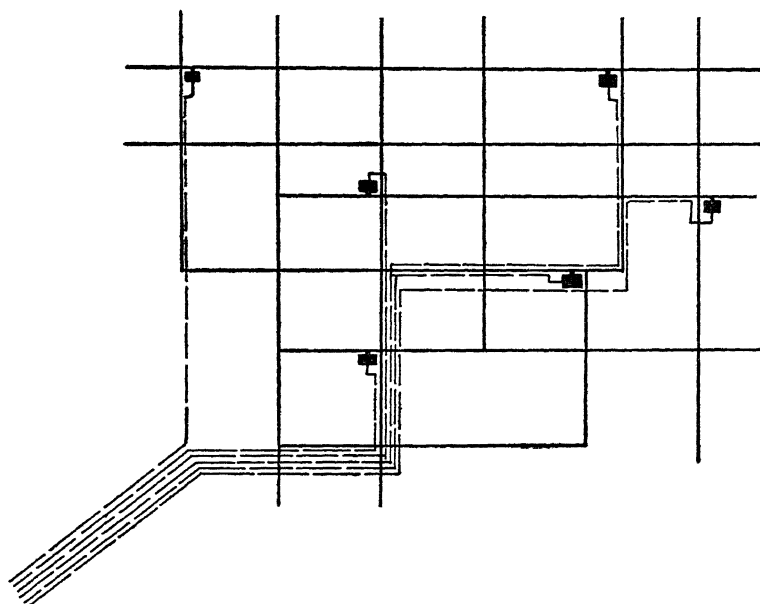


Fig. 4. Alternating Current Distribution with Secondary Mains and Primary Feeders.

are primary feeders of 2200 volts, connecting to the mains through step-down transformers. In such a system, by varying the voltage impressed upon the primary feeders, a voltage regulation of the system similar to that of direct current distribution becomes feasible. Such an arrangement has these advantages over the direct current system: the drop in the feeders is very much lower, due to their higher voltage; and

that the feeder voltage can be regulated by alternating current feeder regulators or compensators, that is, stationary structures similar to the transformer. It has, however, the disadvantage that, due to the self-induction of the mains, each feeding point can supply current over a far shorter distance than with direct current, and the interchange of current between feeders, by which the load can be shifted and apportioned between the feeders, is far less.

As a result, it is difficult to reach as good voltage regulation with the same attention to the system; and since this arrangement has the disadvantage that any breakdown in the secondary system or in a transformer may involve the entire system, this system of inter-connected secondary mains is rarely used for alternating current distribution, but the secondary mains are usually kept separate. That is, as shown diagrammatically in Fig. 5, a number of separate secondary mains are fed by large transformers from primary feeders, and usually each primary feeder connects to a number of transformers. Where the distances are considerable, and the voltage drop in the primary feeders appreciable, voltage regulation of the feeders becomes necessary; and in this case, to get good voltage regulation in the system, attention must be given to the arrangements of the feeders and mains. That is, all the transformers on the same feeder should be at about the same distance from the station, so that the voltage drop between the transformers on the same feeder is negligible; and the nature of the load on the secondary mains fed by the same feeder should be about as nearly the same as feasible, so that all the mains on the same feeder are about equally loaded. It would therefore be undesirable for voltage regulation, to connect, for instance, a main feeding a

residential section to the same feeder as a main feeding a business district or an office building.

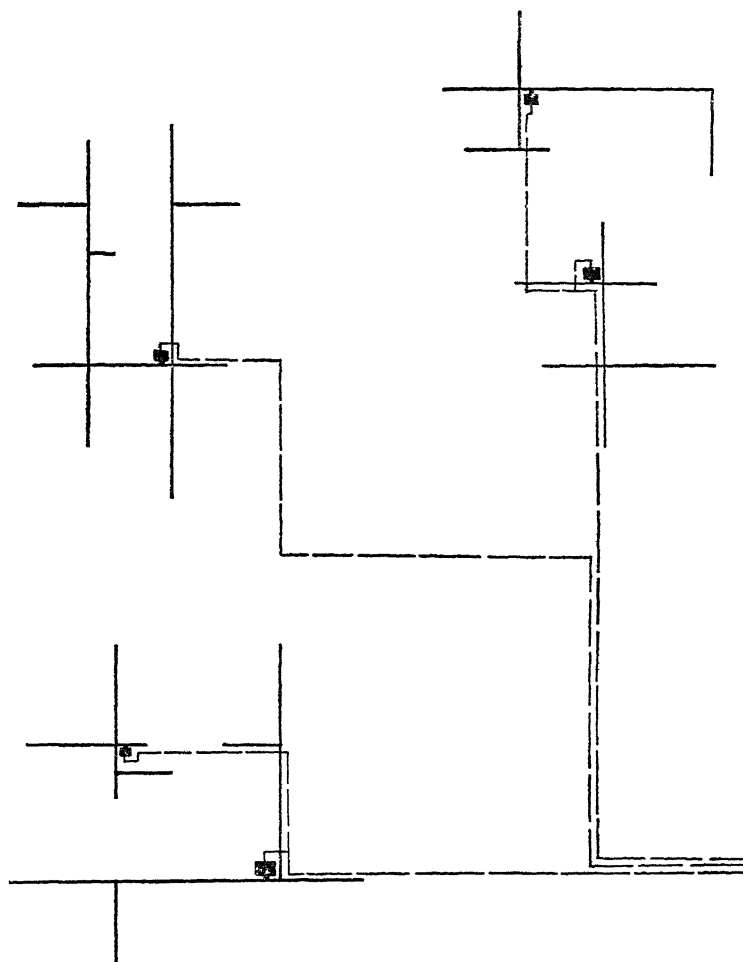


Fig. 5. Typical Alternating Current Distribution.

In a well designed alternating current distribution system, that is, a system using secondary distribution mains as far as feasible, the all year efficiency is about the same as with the direct current system. In such an alternating current system,

the efficiency at heavy load is higher, and at light load lower, than in the direct current system; in this respect the alternating current system has the advantage over the direct current system, since at the time of heavy load the power is more valuable than at light load.

THIRD LECTURE

LIGHT AND POWER DISTRIBUTION

IN A DIRECT current distribution system, the motor load is connected to the outside mains at 220 volts, and only very small motors, as fan motors, between outside mains and neutral; since the latter connection, with a large motor, would locally unbalance a system. The effect of a motor on the system depends upon its size and starting current, and with the large mains and feeders, which are generally used, even the starting of large elevator motors has no appreciable effect, and the supply of power to electric elevators represents a very important use of direct current distribution.

In alternating current distribution systems, the effect on the voltage regulation, when starting a motor, is far more severe; since alternating current motors in starting usually take a larger current than direct current motors starting with the same torque on the same voltage; and the current of the alternating current motor is lagging, the voltage drop caused by it in the reactance is therefore far greater than would be caused by the same current taken by a non-inductive load, as lamps. Furthermore, alternating current supply mains usually are of far smaller capacity, and therefore more affected in voltage. Large motors are therefore rarely connected to the lighting mains of an alternating current system, but separate transformers and frequently separate feeders are used for the motors, and very large motors commonly built for the primary distribution voltage of 2200, are connected to these mains.

For use in an alternating current distribution system, the synchronous motor hardly comes into consideration, since the synchronous type is suitable mainly for large powers, where it is operated on a separate circuit.

The alternating current motor mostly used in small and moderate sizes—such as come into consideration for power distribution from a general supply system—is the induction motor. The single-phase induction motor, however, is so inferior to the polyphase induction motor, that single-phase motors are used only in small sizes; for medium and larger sizes the three-phase or two-phase motor is preferred. This however, introduces a complication in the distribution system, and the three-wire single-phase system therefore is less suited for motor supply, but additional conductors have to be added to give a polyphase power supply to the motor. As the result thereof, motors are not used in alternating current systems to the same extent as in direct current systems. In the alternating current system, however, the motor load is, if anything, more important than in the direct current system, to increase the load factor of the system; since the efficiency of the alternating current system decreases with decrease of load, while that of a direct current system increases.

Compared with the direct current motor, the polyphase induction motor has the disadvantage of being less flexible: its speed cannot be varied economically, as that of a direct current motor by varying the field excitation. Speed variation of the induction motor produced by a rheostat in the armature or secondary circuit, in the so-called form “M” motor is accomplished by wasting power: the power input of an induction motor always corresponds to full speed; if the speed is reduced by running on the rheostat, the difference in power between that which the motor actually gives, and that which it would give, with the same torque, at full speed, is consumed in the rheostat.

Where therefore different motor speeds are required, provisions are made in the induction motor to change the number

of poles; thereby a number of different definite speeds are available, at which the motor operates economically as "multi-speed" motor.

The starting torque of the polyphase induction motor with starting rheostat in the armature (Form L motor) is the same as the running torque at the same current input, just as in the case of the direct current shunt motor with constant field excitation. In the squirrel cage induction motor, however, (Form K motor) the starting torque is far less than the running torque at the same current input; or inversely, to produce the same starting torque, a greater starting current is required. In starting torque or current, the squirrel cage induction motor has the disadvantage against the direct current motor. It has, however, an enormous advantage over it in its greater simplicity and reliability, due to the absence of commutator and brushes, and the use of a squirrel cage armature.

The advantage of simplicity and reliability of the squirrel cage induction motor sufficiently compensates for the disadvantage of the large starting current, to make the motor most commonly used. In an alternating current distribution system, however, great care has to be taken to avoid the use of such larger motors at places where their heavy lagging starting currents may affect the voltage regulation; in such places, separate transformers and even separate primary feeders are desirable.

The single-phase induction motor is not desirable in larger sizes in a distribution system, since its starting current is still larger; in small sizes, however, it is extensively used, since it requires no special conductors, but can be operated from a single-phase lighting main.

The alternating current commutator motor is a single-phase motor which has all the advantages of the different types of direct current motors; it can be built as constant speed motor of the shunt type, or as motor with the characteristics of the direct current series motor: very high starting torque with moderate starting current. It has, however, also the disadvantages of the direct current motor: commutator and brushes; and so requires more attention than the squirrel cage induction motor.

Alternating current generators now are almost always used as polyphase machines, three-phase or two-phase, and transmission lines are always three-phase, though in transforming down, the system can be changed to two-phase. The power supply in an alternating current system therefore is practically always polyphase; and since a motor load, which is very desirable for economical operation, also requires polyphase currents, alternating current distribution systems always start from polyphase power.

The problem of alternating current distribution therefore is to supply, from a polyphase generating system, single-phase current to the incandescent lamps, and polyphase current to the induction motors.

PRIMARY DISTRIBUTION SYSTEMS

1. Two conductors of the three-phase generating or transmission system are used to supply a 2200 single-phase system for lighting by step-down transformers and three-wire secondary mains; the third conductor is carried to those places where motors are used and three-phase motors are operated by separate step-down transformers. In the lighting feeders, the voltage is then controlled by feeder regulators, or, in a smaller system, the generator excitation is varied so as to main-

tain the proper voltage on the lighting phase. At load, the three-phase triangle then more or less unbalances, but induction motors are very little sensitive to unbalancing of the voltage, and by their regulation—by taking more current from the phase of higher, less from the phase of lower voltage—tend to restore the balance. For smaller motors, frequently two transformers are used, arranged in “open delta” connection.

2. Two-phase generators are used, or in the step-down transformers of a three-phase transmission line, the voltage is changed from three-phase to two-phase; the lighting feeders are distributed between the two phases and controlled by potential regulators so that the distribution for lighting is single-phase, by three-wire secondary mains. For motors, both phases are brought together, and the voltage stepped down for use on two-phase motors. This requires four, or at least three, primary wires to motor loads.

3. From three-phase generators or transmission lines, three separate single-phase systems are operated for lighting; that is the lighting feeders are distributed between the three phases, and all three primary wires are brought to the step-down transformers for motors. This arrangement, by distributing the lighting feeders between the three phases, would require more care in exactly balancing the load between all three phases than two, but a much greater unbalancing can be allowed without affecting the voltage.

4. Four-wire three-phase primary distribution with grounded neutral, and 2200 volts between outside conductors and neutral. The lighting feeders are distributed between the three circuits between outside conductors and neutral, and motors supplied by three of such transformers. This system is becoming of increasing importance, since it allows economical distribution to distances beyond those which can be reached

with 2200 volts: with 2600 volts on the transformers—as the upper limit of primary distribution voltage—the voltage between outside conductors is 4500, and the copper economy of the system therefore is that of a 4500 volt three-phase system.

5. Polyphase primary and polyphase secondary distribution, with the motor connected to the same secondary mains as the lights.

SYSTEMS OF LOW TENSION DISTRIBUTION FOR LIGHTING AND POWER.

I. TWO-WIRE DIRECT CURRENT OR SINGLE-PHASE 110 VOLTS. Fig 6.

This can be used only for very short distances, since its copper economy is very low, that is, the amount of conductor material is very high for a given power.

Cu. I.

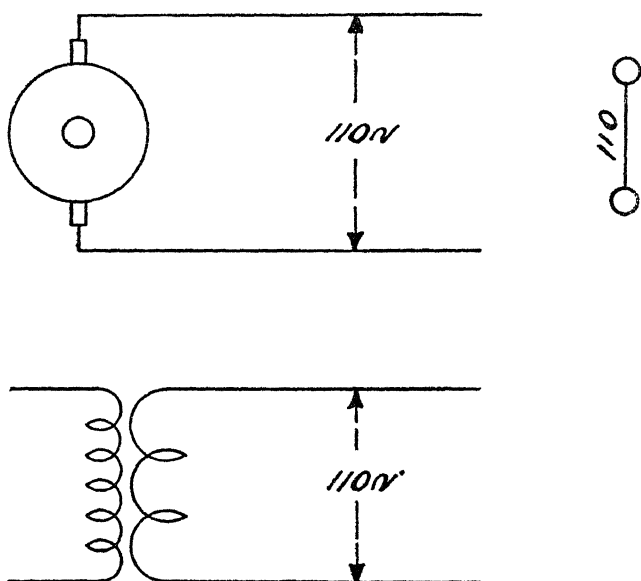


Fig. 6. Two-Wire System.

2. THREE-WIRE DIRECT CURRENT OR SINGLE-PHASE 110-220 VOLTS. Fig. 7.

Neutral one-half size of the two outside conductors. The two outside conductors require one-quarter the copper of the two wires of a 110 volt system; since at twice the voltage and one-half the current, four times the resistance or one-quarter

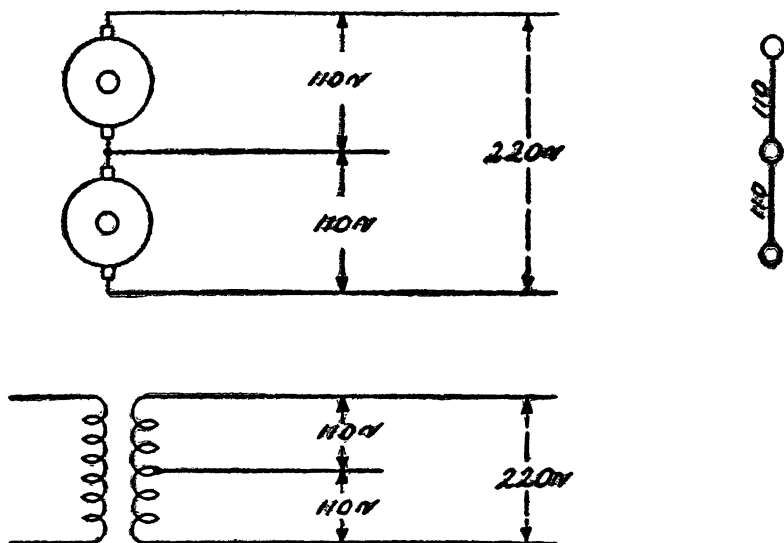


Fig. 7. Three-Wire System.

the copper is sufficient for the same loss (the amount of conductor material varying with the square of the voltage).

Adding then one-quarter for the neutral of half-size, gives $\frac{1}{4} \times \frac{1}{4} = \frac{1}{16}$ or altogether $\frac{1}{4} + \frac{1}{16} = \frac{5}{16}$ of the conductor material required by the two-wire 110 volt system. That is, the copper economy is $\frac{5}{16}$. This is the most commonly used system, since it is very economical, and requires only three conductors. It is, however, a single-phase system, and therefore not suitable for operating polyphase induction motors.

$$\text{Cu. } \frac{5}{16}$$

3. FOUR-WIRE QUARTER-PHASE (TWO-PHASE). Fig. 8.

Two separate two-wire single-phase circuits, therefore no saving in copper over two-wire systems. That is, the copper economy is:

$$\text{Cu. 1.}$$

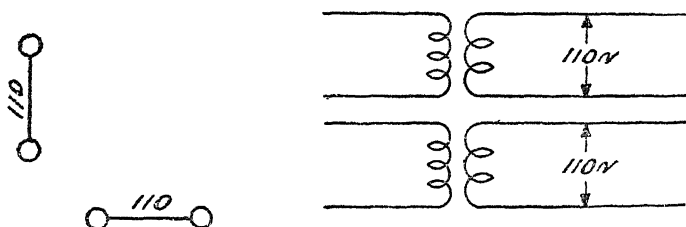


Fig. 8. Four-Wire Two-Phase System.

4. THREE-WIRE QUARTER-PHASE. Fig. 9.

Common return of both phases, therefore saves one wire or one-quarter of the copper; hence has the copper economy:

$$\text{Cu. } \frac{3}{4}$$

In this case however, the middle or common return wire carries $\sqrt{2}$, or 1.41 times as much current as the other two wires, and when making all three wires of the same size, the copper is not used most economically. A small further saving is therefore made by increasing the middle wire and decreasing the

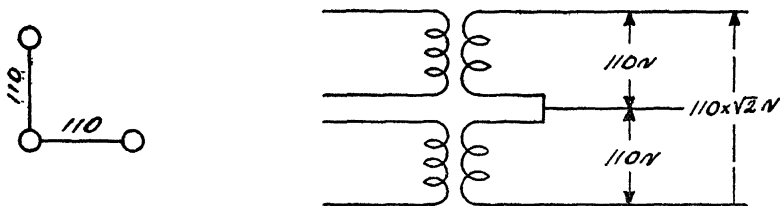


Fig. 9. Three-Wire Two-Phase System.

outside wires so that the middle wire has 1.41 times the section of each outside wire. This improves the copper economy to:

$$\text{Cu. 0.73}$$

5. THREE-WIRE THREE-PHASE. Fig. 10.

A three-phase system is best considered as a combination of three single-phase systems, of the voltage from line to neutral, and with zero return (because the three currents neutralize each other in the neutral).

Compared thereto the two-wire single-phase system can be considered as a combination of two single-phase circuits from wire to neutral with zero return.

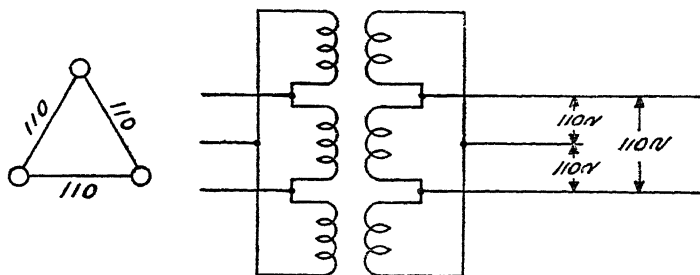


Fig. 10. Three-Wire Three-Phase System.

In a 110 volt single-phase system the voltage from line to neutral equals $\frac{110}{2}$, in a three-phase system equals $\frac{110}{\sqrt{3}}$.

The ratio of voltages is $\frac{110}{2} \div \frac{110}{\sqrt{3}}$, or $\frac{110 \times \sqrt{3}}{2 \times 110} = \frac{\sqrt{3}}{2}$, and the square of the ratio of voltages equals $\frac{3}{4}$; and as the copper economy varies with the square of the voltage, the copper economy for the three-wire three-phase system is:

$$\text{Cu. } \frac{3}{4}$$

6. FIVE-WIRE QUARTER-PHASE. Fig. 11.

Neglecting the neutral conductor, the five-wire quarter-phase system can be considered as four single-phase circuits without return, from line to neutral, of voltage 110. Compared with the two-wire circuit, which consists of two single-phase circuits without return, of $\frac{110}{2}$ volts, No. 6 therefore has twice the voltage of No. 1; therefore one-quarter the copper.

Making the neutral half the size of the main conductor adds one-half of the copper of one conductor, or $\frac{1}{8}$ of $\frac{1}{4} = \frac{1}{32}$, so giving a total of $\frac{1}{4} + \frac{1}{32}$, that is, a copper economy of:

$$\text{Cu.} = \frac{9}{32}$$

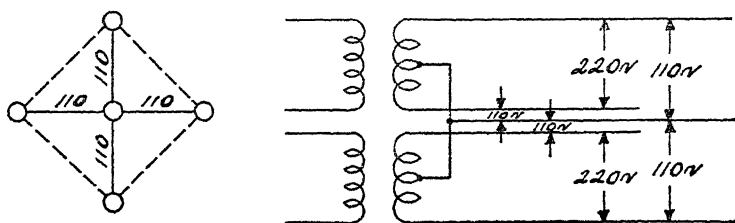


Fig. 11. Five-Wire Two-Phase System.

7. FOUR-WIRE THREE-PHASE. Fig. 12.

Lamps connected between line and neutral.

Neglecting the neutral, the system consists of three single-phase circuits without return, of 110 volts, and compared with

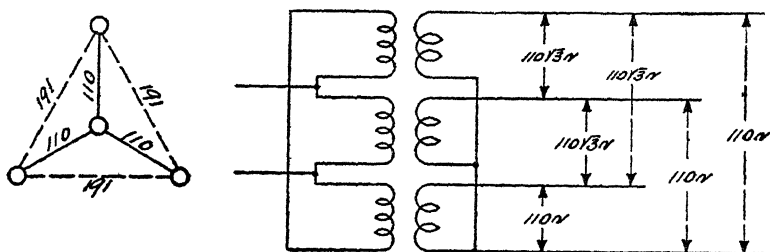


Fig. 12. Four-Wire Three-Phase System.

the two-wire circuit of $\frac{110}{2}$ between wire and neutral without return, it therefore requires one-quarter the copper.

Making the neutral one-half size adds $\frac{1}{6}$ of the copper, or $\frac{1}{6}$ of $\frac{1}{4} = \frac{1}{24}$, and so gives a total copper economy of $\frac{1}{24} + \frac{1}{4} = \frac{7}{24}$.

$$\text{Cu.} = \frac{7}{24}$$

8. THREE-WIRE SINGLE-PHASE LIGHTING WITH THREE-PHASE POWER. Fig. 13.

Lighting: Half size neutral, same as No. 2, therefore copper economy: $\text{Cu.} = \frac{5}{16}$

Power: Three-wire three-phase 220 volts; that is, the same as No. 5, but twice the voltage, thus one-quarter the copper of No. 5, or $\frac{1}{4}$ of $\frac{3}{4} = \frac{3}{16}$: $\text{Cu.} = \frac{3}{16}$

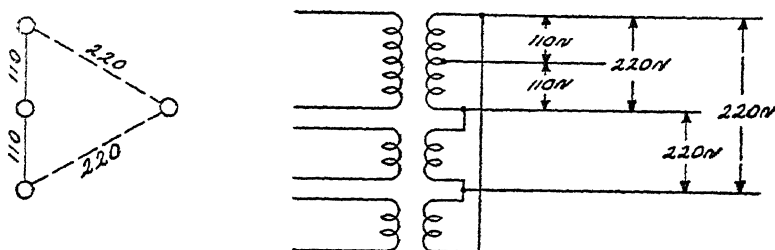


Fig. 13. Single-Phase Lighting and Three-Phase Power.

The systems mostly used are:

No. 2. Three-wire direct current or alternating current single-phase.

No. 8. Three-wire lighting, three-phase power. Less frequent.

No. 6. Five-wire quarter-phase.

No. 7. Four-wire three-phase.

As we have seen, the two-wire system is rather inefficient in copper. High efficiency requires the use of a third conductor, that is, the three-wire system, for direct current or single-phase alternating current.

Three-wire polyphase systems, however, are inefficient in copper, as No. 4 and No. 5; and to reach approximately the same copper economy, as is reached by a three-wire system with direct current and single-phase alternating current, requires at least four wires with a polyphase system.

That is, for equal economy in conductor material, the polyphase system requires at least one more conductor than the single-phase or the direct current distribution system.

While the field of direct current distribution is found in the interior of large cities, alternating current is used in smaller towns and villages and in the suburbs of large cities. In the latter, therefore, alternating current does the pioneer work. That is, the district is developed by alternating current, usually with overhead conductors, and when the load has become sufficiently large to warrant the establishment of converter substations, direct current mains and feeders are laid under ground, the alternating current distribution is abandoned, and the few alternating current motors are replaced by direct current motors. In the last years, however, considerable motor load has been developed in the alternating current suburban distribution systems, fairly satisfactorily alternating current elevator motors have been developed and introduced and the motor load has become so large as to make it economically difficult to replace the alternating current motors by direct current motors in changing the system to direct current; and it therefore appears that the distribution systems of large cities will be forced to maintain alternating current distribution even in districts of such character as would make direct current preferable.

FOURTH LECTURE

LOAD FACTOR AND COST OF POWER

The cost of the power supplied at the customer's meter consists of three parts.

A. A fixed cost, that is, cost which is independent of the amount of power used, or the same whether the system is fully loaded or carries practically no load. Of this character, for instance, is the interest on the investment in the plant, the salaries of its officers, etc.

B. A cost which is proportional to the amount of power used. Such a proportional cost, for instance, is that of fuel in a steam plant.

C. A cost depending on the reliability of service required, as the cost of keeping a steam reserve in a water power transmission, or a storage battery reserve in a direct current distribution.

Since of the three parts of the cost, only one, B, is proportional to the power used, hence constant per kilowatt output,—the other two parts being independent of the output,—hence the higher per kilowatt, the smaller a part of the capacity of the plant the output is; it follows that the cost of power delivered is a function of the ratio of the actual output of the plant, to the available capacity.

Interest on the investment of developing the water power or building the steam plant, the transmission lines, cables and distribution circuits, and depreciation are items of the character A, or fixed cost, since they are practically independent of the power which is produced and utilized.

Fuel in a steam plant, oil, etc., are proportional costs, that is, essentially depending on the amount of power produced.

Salaries are fixed cost, A; labor, attendance and inspection are partly fixed cost A, partly proportional cost B,—economy of operation requires therefore a shifting of as large a part thereof over into class B, by shutting down smaller substations during periods of light load, etc.

Incandescent lamp renewals, arc lamp trimming, etc., are essentially proportional costs, B.

The reserve capacity of a plant, the steam reserve maintained at the receiving end of a transmission line, the difference in cost between a duplicate pole line and a single pole line with two circuits, the storage battery reserve of the distribution system, the tie feeders between stations, etc., are items of the character C; that is, part of the cost insuring the reliability and continuity of power supply.

The greater the fixed cost A is, compared with the proportional cost B, the more rapidly the cost of power per kilowatt output increases with decreasing load. In steam plants very frequently A is larger than B, that is, fuel, etc. not being the largest items of cost; in water power plants A practically always is far larger than B. As result thereof, while water power may appear very cheap when considering only the proportional cost B—which is very low in most water powers—the fixed cost A usually is very high, due to the hydraulic development required. The difference in the cost of water power from that of steam power therefore is far less than appears at first. As water power is usually transmitted over a long distance line, while steam power is generated near the place of consumption, water power usually is far less reliable than steam power. To insure equal reliability, a water power plant brings the item C, the reliability cost, very high in comparison with the reliability cost of a steam power plant, since the possibility of a break-down of a transmission line requires a steam reserve, and

where absolute continuity of service is required, it requires also a storage battery, etc. ; so that on the basis of equal reliability of service, sometimes very little difference in cost exists between steam power and water power, unless the hydraulic development of the latter was very simple.

The cost of electric power of different systems therefore is not directly comparable without taking into consideration the reliability of service and the character of the load.

As a very large, and frequently even the largest part of the cost of power, is independent of the power utilized, and therefore rapidly increases with decreasing load on the system, the ratio of average power output to the available power capacity of the plant is of fundamental importance in the cost of power per kilowatt delivered. This ratio, of the average power consumption to the available power, or station capacity, has occasionally been called "load factor." This definition of the term "load factor" is, however, undesirable, since it does not take into consideration the surplus capacity of the station, which may have been provided for future extension; the reserve for insuring reliability C, etc.; and other such features which have no direct relation whatever to the character of the load.

Therefore as load factor is understood, in accordance with the definition in the Standardization Rules of the A. I. E. E., the ratio of the average load to the maximum load; any excess of the station capacity beyond the maximum load is power which has not yet been sold, but which is still available for the market, or which is held in reserve for emergencies, is not charged against the load factor.

The cost of electric power essentially depends on the load factor. The higher the load factor, the less is the cost of the power, and a low load factor means an abnormally high cost

per kilowatt. This is the case in steam power, and to a still greater extent in water power.

For the economical operation of a system, it therefore is of greatest importance to secure as high a load factor as possible, and consequently, the cost—and depending thereon the price—of electric power for different uses must be different if the load factors are different, and the higher the cost, the lower the load factor.

Electrochemical work gives the highest load factor, frequently some 90%, while a lighting system shows the poorest load factor—in an alternating current system without motor load occasionally it is as low as 10 to 20%.

Defining the load factor as the ratio of the average to the maximum load, it is necessary to state over how long a time the average is extended; that is, whether daily, monthly or yearly load factor.

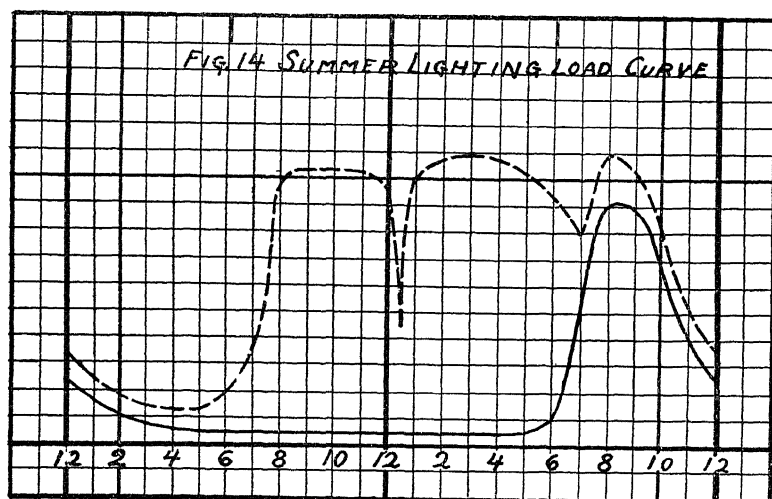


Fig. 14. Summer Lighting Load Curve.

For instance, Fig. 14 shows an approximate load curve of a lighting circuit during a summer day: practically no load

except for a short time during the evening, where a high peak is reached. The ratio of the average load to the maximum load during this day, or the daily load factor, is 22.8%.

Fig. 15 shows an approximate lighting load curve for a winter day: a small maximum in the morning, and a very high evening maximum, of far greater width than the summer day curve, giving a daily load factor of 34.5%.

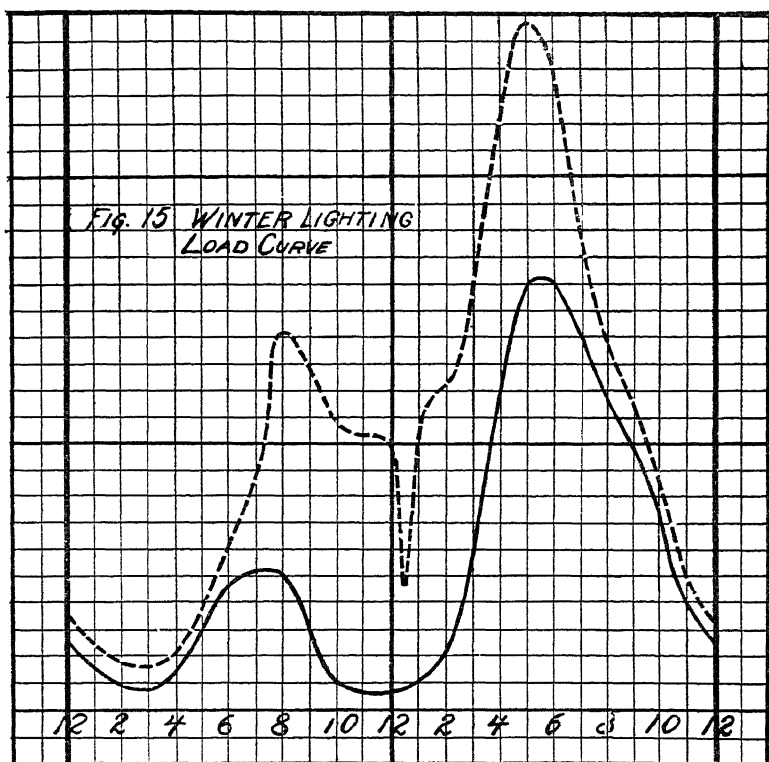


Fig. 15. Winter Lighting Load Curve.

During the year, the daily load curve varies between the extremes represented by Figs. 14 and 15, and the average annual load is therefore about midway between the average load of a summer day and that of a winter day. The maximum yearly load, however, is the maximum load during the winter

day; and the ratio of average yearly load to maximum yearly load, or the yearly load factor of the lighting system, therefore is far lower than the daily load factor: if we consider the average yearly load as the average between 14 and 15, the yearly load factor is only 23.6%.

One of the greatest disadvantages of lighting distribution therefore is the low yearly load factor, resulting from the summer load being so very far below the winter load; economy of operation therefore makes an increase of the summer lighting load very desirable. This has led to the development of spectacular lighting during the summer months, as represented by the various Luna Parks, Dreamlands, etc.

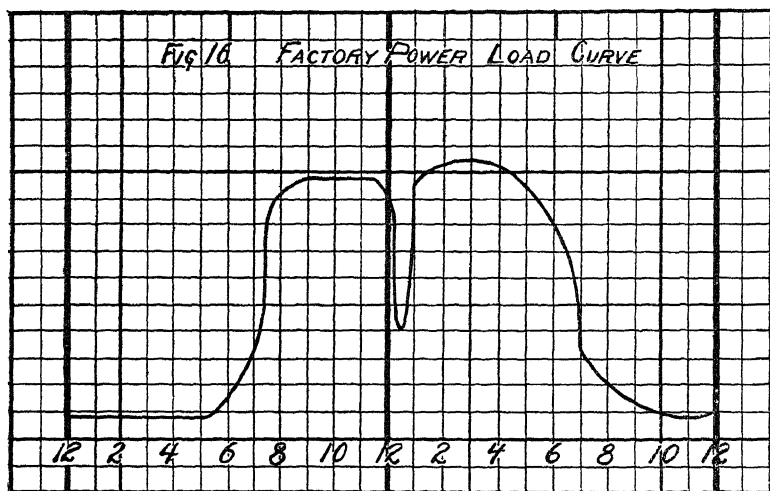


Fig. 16. Factory Power Load Curve.

The load curve of a factory motor load is about the shape shown in Fig. 16: fairly constant from the opening of the factories in the morning to their closing in the evening, with perhaps a drop of short duration during the noon hour, and a low extension in the evening, representing overtime work. It gives a daily load factor of 49.5%.

This load curve, superimposed upon the summer lighting curves, does not appreciably increase the maximum, but very greatly increases the average load, as shown by the dotted curve in Fig. 14; and so improves the load factor, to 65.4%—thereby greatly reducing the cost of the power to the station, in this way showing the great importance of securing a large motor load. During the winter months, however, the motor load overlaps the lighting maximum, as shown by the dotted curve in Fig. 15. This increases the maximum, and thereby increases the load factor less, only to 41.7%. This is not so serious in the direct current system with storage battery reserve, as the overlap extends only for a short time, the overload being taken care of by storage batteries or by the overload capacity of generators and steam boilers; but where it is feasible, it is a great advantage if the users of motors can be induced to shut them down in winter with beginning darkness.

It follows herefrom, that additional load on the station during the peak of the load curve is very expensive, since it increases the fixed cost A and C, while additional load during the periods of light station load, only increases the proportional cost B; it therefore is desirable to discriminate against peak loads in favor of day loads and night loads. For this purpose, two-rate meters have been developed, that is, meters which charge a higher price for power consumed during the peak of the load curve, than for power consumed during the light station loads. To even out load curves, and cut down the peak load, maximum demand meters have been developed, that is, meters which charge for power somewhat in proportion to the load factor of the circuit controlled by the meter. Where the circuit is a lighting circuit, and the maximum demand therefore coincides with the station peak,

this is effective, but on other classes of load the maximum demand meters may discriminate against the station. For instance, a motor load giving a high maximum during some part of the day, and no load during the station peak, would be preferable to the station to a uniform load throughout the day, including the station peak, while the maximum demand meter would discriminate against the former.

By a careful development of summer lighting loads and motor day loads, the load factors of direct current distribution systems have been raised to very high values, 50 to 60% ; but in the average alternating current system, the failure of developing a motor load frequently results in very unsatisfactory yearly load factors.

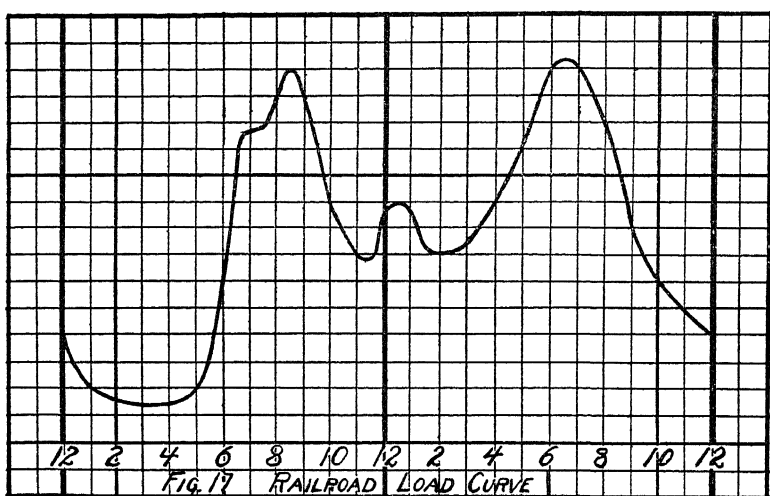


Fig. 17. Railroad Load Curve.

The load curve of a railway circuit is about the shape of that shown in Fig. 17 : a fairly steady load during the day, with a morning peak and an evening peak, occasionally a smaller noon peak and a small second peak later in the evening, then tapering down to a low value during the night. The average

load factor usually is far higher than in a lighting circuit, in Fig. 17: 54.3%.

In defining the load factor, it is necessary to state not only the time over which the load is to be averaged, as a day, or a year, but also the length of time which the maximum load must last, must be counted. For instance, a short circuit of a large motor during peak load, which is opened by the blowing of the fuses, may momentarily carry the load far beyond the station peak without being objectional. The minimum duration of maximum load, which is chosen in determining the load factor, is that which is permissible without being objectionable for the purpose for which the power is distributed. Thus in a lighting system, where voltage regulation is of foremost importance, minutes may be chosen, and maximum load may be defined as the average load during that minute during which the load is a maximum; while in a railway system, a half-hour may be used as a duration of maximum load, as a railway system is not so much affected by a drop of voltage due to overload, and an overload of less than half an hour may be carried by the overload capacity of the generators and the heat storage of the steam boilers; so that a peak load requires serious consideration only when it exceeds half an hour.

FIFTH LECTURE

LONG DISTANCE TRANSMISSION

THREE-PHASE is used altogether for long distance transmission. Two-phase is not used any more, and direct current is being proposed, having been used abroad in a few cases: but due to the difficulty of generation and utilization, it is not probable that it will find any extended use, so that it does not need to be considered.

FREQUENCY

The frequency depends to a great extent on the character of the load, that is, whether the power is used for alternating current distribution—60 cycles—or for conversion to direct current—25 cycles. For the transmission line, 25 cycles has the advantage that the charging current is less and the inductive drop is less, because charging current and inductance voltage are proportional to the frequency.

VOLTAGE

11,000 to 13,200 volts and more recently, even 22,000 volts is most common for shorter distances, as 10 to 20 miles, since this is about the highest voltage for which generators can be built; its use therefore saves the step-up transformers, that is, the generator feeds directly into the line and to the step-down transformers for the regular load.

The next step is 30,000 volts; that is, 33,000 volts at the generator, 30,000 at the receiving end of the line. No intermediate voltages between this and the voltage for which generators can be wound is used, as 30,000 volts does not yet offer any insulator troubles; but line insulators can be built at moderate cost for this voltage, and as step-up transformers

have to be used, it is not worth while to consider any lower voltage than 33,000 volts. This voltage transmits economically up to distances of 50 to 60 miles.

40,000 to 44,000 volts is the next step; it is used for high power transmission lines of greater distance, where reliability of operation is of importance and the use of a conservative voltage therefore preferable to the attempt at economizing by the use of extra high voltages.

A number of 60,000 volt systems are in more or less successful operation, and systems of 80,000 to 110,000 volts are in construction and a few in operation. Where the distances are very great, power valuable, and continuity of service not of such foremost importance, such voltages are justified in the present state of the art.

In such very high voltage systems, the transformers are occasionally wound so that they can be connected for half voltage, for operating the line at half voltage, until the load has sufficiently increased to require full voltage; or the transformers are built for star or Y connection at full voltage, and at first operated in ring or delta connection, at $\frac{1}{\sqrt{3}} = 57\%$ of full voltage.

The cost of a long distance transmission line depends on the voltage used.

The cost of line conductors decreases with the square of the voltage.

At twice the voltage, twice the line drop can be allowed with the same loss; at twice the voltage the current is only half for the same power, and twice the drop with half the current gives four times the resistance, that is, one-quarter the conductor section and cost.

The cost of line insulators increases with increase of voltage. The cost of pole line increases with increase of voltage, since greater distance between the conductors is necessary and so longer poles, longer cross arms, and heavier construction, and not so many circuits can be carried on the same pole line.

The lower the voltage, the greater in general is the reliability of operation, since a larger margin of safety can be allowed.

Since a part of the cost of the transmission line decreases, another part increases with the voltage, a certain voltage will be most economical.

Lower voltage increases the cost of the conductor, higher voltage increases the cost of insulators and line construction, and decreases the reliability.

The most economical voltage of a transmission line varies with the cost of copper. When copper is very high, higher voltages are more economical than when copper is low. The same applies to aluminum, since the price of aluminum has been varied with that of copper.

Aluminum generally is used as stranded conductor. In the early days single wire gave much trouble by flaws in the wire. Aluminum expands more than copper with temperature changes, and so when installing the line in summer, a greater sag must be allowed than with copper, otherwise it stretches so tight in winter that it may tear apart. Aluminum also is more difficult to join together, since it cannot be welded.

For the same conductivity an aluminum line has about twice the size, but one-half of the weight of a copper conductor, and costs 10% less; but copper has a permanent value, while the price of aluminum may sometime drop altogether, as the metal has no intrinsic value, being one of the most common

constituents of the surface of the earth, and its cost is merely that of its separation or reduction.

LOSSES IN LINE DUE TO HIGH VOLTAGE

The loss in the line by brush discharge or corona effect is nothing up to a certain voltage, but at a certain voltage it begins and very rapidly increases.

The voltage at which a loss by corona effect begins is where the air at the surface of the conductor breaks down, becomes conducting and thus luminous. This occurs at a potential gradient of 100,000 to 120,000 volts per inch.

The potential gradient is highest at the surface of the conductor.

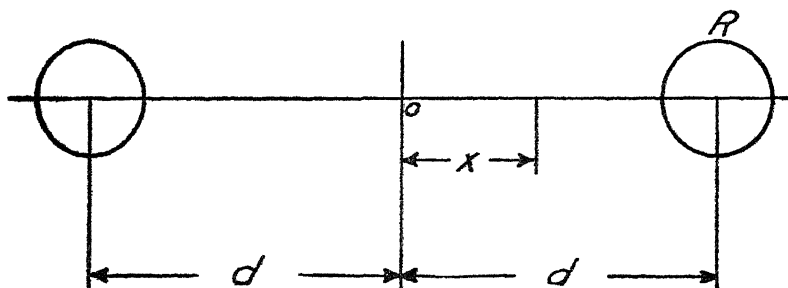


Fig. 18.

In Fig. 18 let

R = radius of conductor.

$2d$ = distance between conductor centres.

At a point x from the centre O the potential is:

$$Q = \frac{c}{d - x} - \frac{c}{d + x} = \frac{2cx}{d^2 - x^2}$$

for: $x = d - R$

that is, at the conductor surface, it is:

$$Q_1 = e$$

Substituting this in the equation, gives :

$$e = \frac{c}{R}$$

hence :

$$c = e R$$

therefore the potential at point x is :

$$f = \frac{2 R x}{d^2 - x^2} e$$

and the potential gradient :

$$g = \frac{d f}{d x} = \frac{2 R (d^2 + x^2)}{(d^2 - x^2)^2} e$$

hence for : $x = d - R$ or the conductor surface : $g_0 = \frac{c}{R}$

If this potential gradient becomes greater than the breakdown strength of air, or 100,000 volts per inch, corona effects and energy losses take place :

$$\frac{e}{R} = 100,000$$

gives :

$e = 100,000 R$ or $E = 100,000 D$, as the voltage where the corona begins, and :

$$R = \frac{e}{100,000} \text{ or } D = \frac{E}{100,000} \text{ is the smallest radius}$$

which can be used, at voltage E , where D is the conductor diameter $= 2 R$, and E is the voltage between the conductors $= 2 e$.

For instance, wire No. 0000

$D = .46''$; corona effects begin at the voltage $E = 100,000 D = 46,000$.

For 100,000 volts the smallest diameter for which no corona effects occur is :

$$D = \frac{E}{100,000} = 1''$$

In high potential transformers in the coils no corona effects may occur, because the diameter of the coil or the thickness is large enough, but the leads connecting the coils with each other and with the outside, if not chosen very large in diameter, may give corona effects and so break down.

In a line or transformer, if one side is grounded, the other side has full voltage against ground, and so may give corona effects and break down; while if not grounded, both sides have half voltage against ground and so give no corona effect. In the first case, the line or transformer so may break down, although the potential differences between the terminals are no greater than in the second case.

For instance, in a 100,000 volt transformer or line, from each terminal to ground are 50,000 volts, and if the conductor diameter is $\frac{1}{2}$ ", no corona effects occur. If now one terminal is grounded, the other terminal has 100,000 volts to ground and so at $\frac{1}{2}$ " diameter gives corona effects, that is, glow and streamers which may destroy the insulating material or produce high frequency oscillations.

At very high voltages it is therefore necessary to have the system statically balanced or symmetrical, that is, have the same potential differences from all the conductors to the ground.

Any electric circuit, and so also the transmission line, contains inductance and capacity, and therefore stores energy as electromagnetic energy in the magnetic field due to the current, and as electrostatic energy, or electrostatic charge, due to the voltage.

If:

e = voltage, C = capacity.

i = current, L = inductance.

the electrostatic energy is:

$$\frac{e^2 C}{2}$$

and the electromagnetic energy:

$$\frac{i^2 L}{2}$$

In a high potential transmission line both energies are of about the same magnitude, and the energy can therefore see-saw between the two forms and thereby produce oscillations and surges resulting in the production of high voltages, which are not liable to occur in circuits in which one of the forms of stored energy is small compared with the other.

In distribution systems up to 2200 volts and even somewhat higher, the electrostatic energy is still negligible and only the electromagnetic energy appreciable.

In static machines the electrostatic energy is appreciable, but the electromagnetic energy negligible.

LINES AND TRANSFORMERS

At voltages above 25,000 step-up and step-down transformers are always used, which are therefore a part of the high potential circuit.

Three-phase is always used in the transmission line.

Some of the available transformer connections are given in Figs. 19 and 20.

Grounding the neutral of the system has the advantage of maintaining static balance and so avoiding oscillations and disturbances in case of an accidental static unbalancing, as for

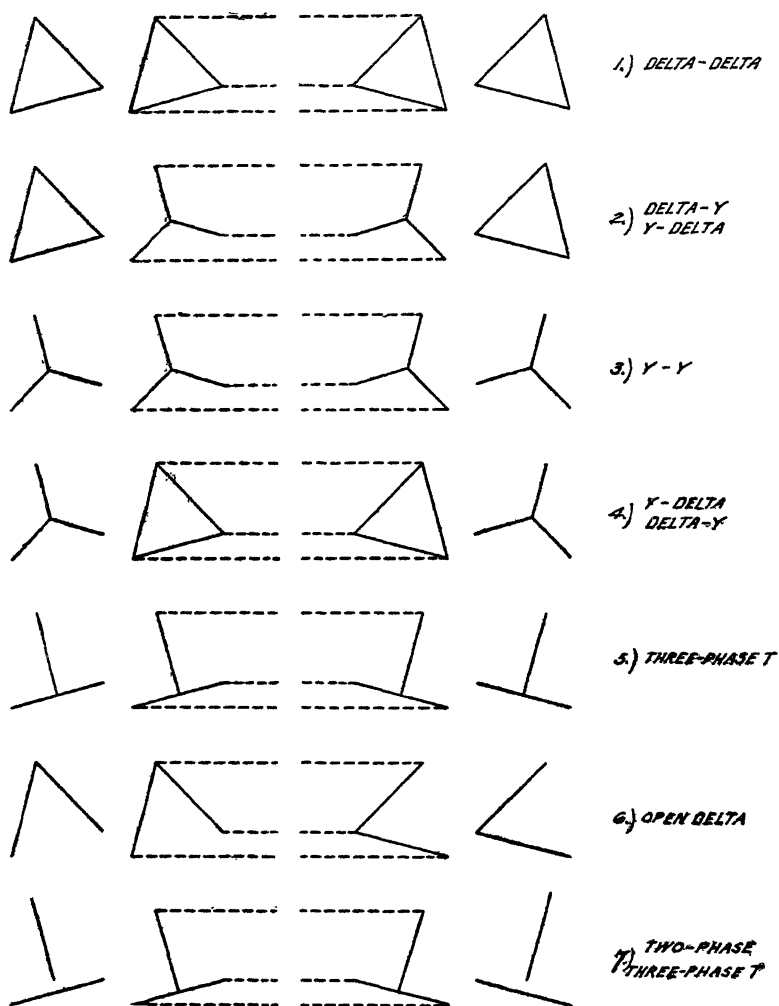
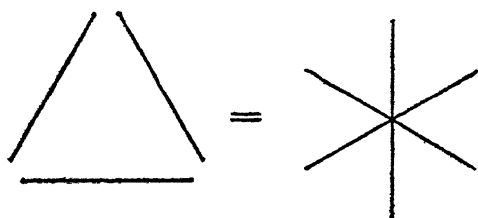


Fig. 19. Transformer Connections.

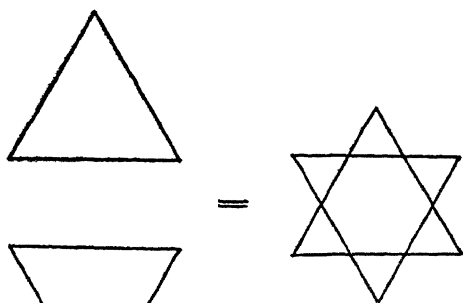
instance, the grounding of one line. It has the disadvantage that a ground on one circuit is a short circuit and so shuts down the circuit.

In connections 1, 4 and 6 no neutral is available for grounding and so three separate transformers have to be installed in Y connection for getting the neutral.

In connections 2 and 3 the neutral can be brought out from the transformer neutral.



8. SIX PHASE DIAMETRICAL



9 SIX PHASE DOUBLE DELTA

Fig. 20. Six-Phase Transformer Connections.

In the T connection 5 and 7, the neutral is brought out from a point at one-third of the teaser transformer winding.

Assuming the line properly installed and insulated, breakdowns may occur, either from mechanical accidents or by high voltages appearing in the line.

HIGH VOLTAGE DISTURBANCES IN TRANSMISSION LINES

These may be:

A. Of fundamental frequency, that is, the same frequency as the alternating current machine circuit.

B. Some higher harmonic of the generator wave, that is, some odd multiple of the generator frequency.

C. Of frequencies entirely independent of the generator, or of a frequency which originates in the circuit, that is, high frequency oscillations as arcing grounds, etc.

If a capacity is in series with an inductance, as the line capacity and the line inductance, the capacity reactance and the inductive reactance are opposed to each other; if they happened to be equal they would neutralize each other, the current would depend on the resistance only and therefore be very large, and with this very large current passing through the inductance and capacity, the voltage at the inductance and at the capacity would be very high.

For instance, if we have 20,000 volts supplied to a circuit having a resistance of 10 ohms and a capacity reactance of 1000 ohms, then the total impedance of the circuit is $\sqrt{10^2 + 1000^2} = 1000$ and the current in the circuit $\frac{20,000}{1000} = 20$ amperes.

If now in addition to the 10 ohms resistance and 1000 ohms capacity reactance, the circuit contains 1000 ohms inductive reactance, the total reactance of the circuit is $1000 - 1000 = 0$ ohms, and the impedance is the same as the resistance, or 10 ohms. The current therefore $\frac{e}{z} = \frac{e}{r} =$

2000 amperes, and the voltage at the capacity therefore is: capacity reactance times amperes = 2,000,000 volts, and the same voltage exists at the inductive reactance.

These voltages are far beyond destruction. That is, if in a circuit of low resistance and high capacity reactance, a high inductive reactance is put in series with the capacity reactance, excessive voltages are produced.

In a transmission line the capacity of the line consumes for instance 10% of full load current; that is, full load voltage sends only 10% of full load current through the capacity. To send full load current through the capacity so would require 10 times full load voltage.

With a line reactance of 20%, 20% or $\frac{1}{5}$ of full load voltage sends full load current through the inductive reactance, while 10 times full load voltage is required by the capacity reactance; the capacity reactance therefore is about 50 times and therefore cannot build up with it to excessive voltages; but to get resonance with the fundamental frequency requires an inductive reactance about 50 times greater than the line reactance.

The only reactance in the system which is large enough to build up with the capacity reactance is the open circuit reactance of the transformers. This is of about the same size as the capacity reactance, since a transformer at open circuit and full voltage takes about 10% of full load current, and the capacity reactance also takes about 10% of full load current.

If therefore a high potential coil of a transformer at open secondary circuit is connected in series with a transmission line, destructive voltages may be produced, by the reactance of the transformer building up with the line capacity. In those transformer connections in which several high

potential coils of different transformers are connected between the transmission wires, this may occur if the low tension coil of one of the transformers accidentally opens and the high potential coil of this transformer then acts as inductive reactance in series with the line capacity in the circuit of the other transformer.

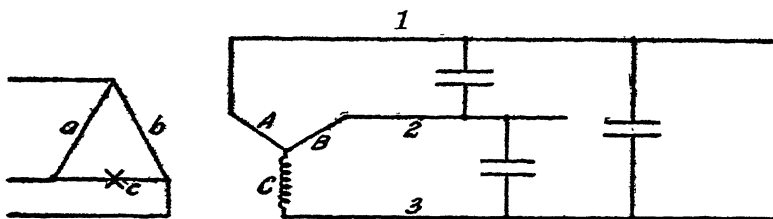


Fig. 21.

This may occur for instance in transformer connection 2, Fig. 19, if as shown in Fig. 21, the low tension coil *c* opens. Then the high tension coil *C* is an inductive reactance in series

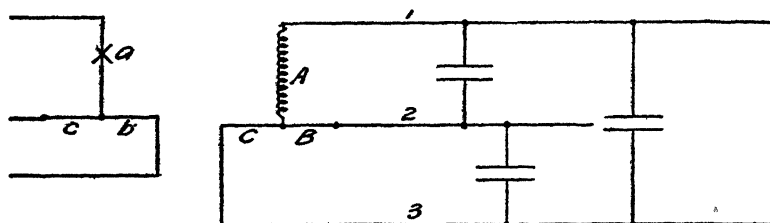


Fig. 22.

with the line capacity from 3 to 1, energized by transformer *A*; and *C* is a high inductive reactance in series with the line capacity from 3 to 2 in a circuit of voltage *B*. That is, from 3 to 1 and from 3 to 2 excessive voltages are produced. So also in T connection, Fig. 22, if for instance the low tension coil *a* opens, the corresponding high tension coil *A* is a high inductive reactance in series with the line capacities in a circuit

of the voltages of the two halves, B and C, of the other transformer, and excessive voltages therefore appear from 1 to 2 and from 1 to 3.

This danger of excessive voltages by the accidental opening of a transformer low tension coil does not exist in delta connection, since in this always only one transformer connects from line to line. It is greatly reduced since the use of triple pole switches became general; and is very much less where several sets of transformers are used in multiple, since even if in one set a low tension coil opens, the other sets maintain the voltage triangle.

Especially dangerous in this respect therefore is the L connection No. 6; since in this case, when using two transformers in open delta, for smaller systems only one set is installed and an accident to one of the transformers causes excessive voltages between its line and the two other lines.

The open circuit reactance of the transformer is the only reactance high enough to give destructive voltages at generator frequency, and in high potential disturbances, the transformer connections should first be carefully investigated to see whether this has occurred.

SIXTH LECTURE

HIGHER HARMONICS OF THE GENERATOR WAVE

THE open circuit reactance of the transformer is the only reactance high enough to give resonance with the line capacity at fundamental frequency.

All other reactances are too low for this.

Since, however, the inductive reactance increases and the capacity reactance decreases proportionally to the frequency, the two reactances come nearer together for higher frequency; that is, for the higher harmonics of the generator wave, and for some of the higher harmonics of the generator wave resonance rise of voltage so may occur between the line capacity and the circuit inductance.

The origin and existence of higher harmonics therefore bears investigation in transformers, transmission lines and cable systems.

ORIGIN OF HIGHER HARMONICS

Higher harmonics may originate in synchronous machines, as generators, synchronous motors and converters, and in transformers.

These two classes of higher harmonics are very different. The former have constant potential character; the latter, constant current character; their cure and prevention therefore must be different, and the method of elimination of one may be very harmful with the other type of harmonics. For instance, the voltage produced by a constant current harmonic as coming from a transformer is eliminated by short circuit. Short circuiting a generator harmonic, however, gives large

short circuit currents, due to the constant potential character, and is therefore dangerous.

HIGHER HARMONICS OF SYNCHRONOUS MACHINES

In synchronous machines, as alternating current generators, the higher harmonics are:

AT NO LOAD

1st. The distribution of magnetism in the air gap depends on the shape of the field poles; it is not a sine wave; neither is the e. m. f. induced by it in an armature a sine wave.

Since there are a number of conductors in series on the armature, the voltage wave is more evened out than that of a single conductor; but still it is not a sine wave, that is, contains harmonics of which the third is the lowest.

2nd. The change of magnetic flux by the passage of open armature slots over the field pole produces harmonics of e. m. f.; that is, when a large open armature slot stands in front of the field pole, the magnetic reluctance is high; the magnetism is lower than when no slot is in front of the field pole; that is, by the passage of the armature slots the field magnetism pulsates, the more so the larger the slots and the fewer they are.

If there are n slots per pole, this produces the two harmonics $2n - 1$ and $2n + 1$.

AT LOAD

3rd. The armature reaction of a single-phase machine pulsates between zero at zero current and a maximum at maximum current.

The resultant armature reaction of a polyphase machine is constant, but locally there is a pulsation making as many cycles per pole as there are phases.

Since the field magnetism under load is due to the combination of field excitation and armature reaction, the pulsation of armature reaction therefore causes a pulsation of field magnetism, and thereby higher harmonics of the e. m. f. wave.

If m = number of phases, the higher harmonics: $2m - 1$ and $2m + 1$ are produced.

4th. The terminal voltage under load is the resultant of the induced e. m. f. and the e. m. f. consumed by the reactance of the armature circuit; that is, the reactance produced by the magnetic flux produced by the armature current in the armature iron. This armature reactance is not constant, but periodically varies, more or less, with double frequency; that is, when the armature coil is in front of the field pole its magnetic circuit is different than when it is between the field poles, and the reactance therefore is different.

This pulsation of armature reactance produces the third harmonic, since it is of double frequency.

The most common and prominent harmonic so is the third harmonic in a synchronous machine.

These harmonics of synchronous machines are induced e. m. f.'s, that is, constant potential or approximately so.

HIGHER HARMONICS OF TRANSFORMERS

In a transformer the wave of e. m. f. depends on that of the magnetism and vice versa. That is, with a sine wave of e. m. f., the magnetism must also be a sine wave, and if the magnetism is not a sine wave, but contains higher harmonics, the e. m. f. is not a sine wave, but contains the harmonics induced by the harmonics of magnetism.

The exciting current of the transformer depends on the magnetism by the hysteresis cycle; if the magnetism is a sine wave, the exciting current therefore cannot be a sine wave, but

must contain higher harmonics—mainly the third harmonic, which reaches 20 to 30% of the fundamental, or even more at saturation.

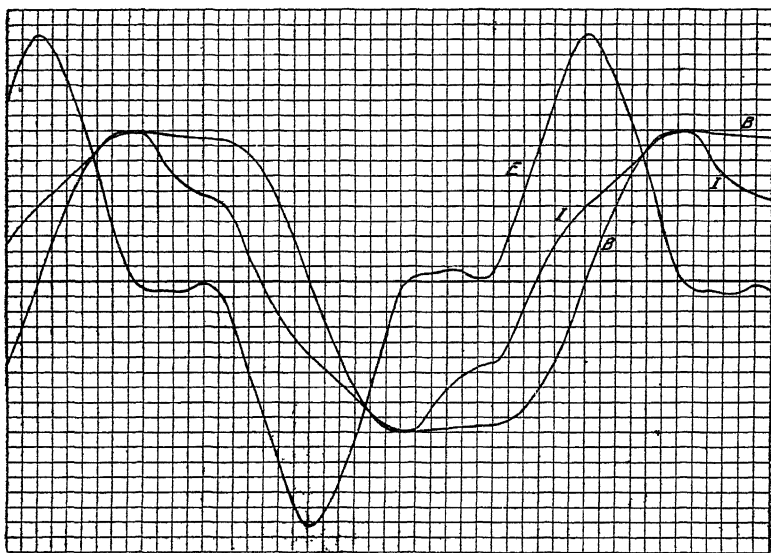


Fig. 23.

In a transformer, e. m. f. and exciting current therefore cannot both be sine waves, but a sine wave of e. m. f. requires an exciting current containing a third harmonic; and a sine wave of exciting current in a transformer or reactive coil thus produces a third harmonic of e. m. f.

If therefore in a transformer the third harmonic is suppressed, and if this third harmonic should have been 20% of the fundamental, then its suppression produces a third harmonic of magnetism of 20% in the opposite direction. A third harmonic of magnetism, however, of 20%, induces a third harmonic of e. m. f. of $3 \times 20 = 60\%$; the e. m. f. being proportional to magnetism and frequency.

The third harmonic of exciting current is positive at the maximum of magnetism, and the third harmonic of magnetism is negative at the maximum, hence is zero and rising at the zero of the magnetism; and at this moment the e. m. f. induced by the third harmonic and by the fundamental therefore are both maxima and in the same direction, that is, add. The suppression of the third harmonic of exciting current thus produces a very high third harmonic of e. m. f., which greatly increases the maximum e. m. f.; that is, the e. m. f. wave is very low for a large part of the cycle and then rises to a very high peak, as shown by Fig. 23; and the maximum e. m. f. may exceed that of a sine wave by 50% and more, thus giving high insulation stress and the possibility of resonance voltages.

EFFECTS OF HIGHER HARMONICS

In a three-phase system the three phases are 120° apart, and their third harmonics are $3 \times 120^\circ = 360^\circ$ apart, that is, in phase with each, and for the third harmonic the three-phase system therefore is a single-phase system.

In a balanced three-phase system, the third harmonics can not exist in the voltages between the lines and in the line currents, if there is no return over the neutral. The three voltages between lines, from 1 to 2, 2 to 3, and 3 to 1, must add up to zero; but since the third harmonics would be in phase with each other, they would not add up to zero, therefore they cannot exist. The three currents, if there is no return over the neutral or the ground, must add up to zero; and since their third harmonics must be in phase with each other, they must be absent. In a balanced three-phase system, third harmonics can exist only in the voltage from line to neutral or Y voltage, in the current from line to line or delta current, and in the

line current only if there is a neutral return or ground return to the generator neutral or transformer neutral.

In a three-phase generator, if the e. m. f. of one phase contains a third harmonic, as is usually the case, then by connecting the three phases in delta connection, the third harmonics of the generator e. m. f.'s are short circuited and so produce a triple frequency current circulating in the generator delta. This triple frequency circulating current can be measured by connecting an ammeter in one corner of the generator delta, and the sum of voltages of the three third harmonics can be measured by putting a voltmeter in a corner of the generator delta. This local current in the generator winding is the triple frequency voltage divided by the generator impedance (the stationary impedance, at triple frequency, but not the synchronous impedance, since the latter includes armature reaction). In generators of low impedance or close regulation, as turbine alternators, this local current may be far more than full load current; delta connection of generator windings therefore is unsafe. As a result, generator windings are almost always connected in Y. Even with delta connection of generator windings no triple frequency appears at the terminals, since its voltage disappears by short circuit.

If the generator winding is connected in Y, the triple frequency voltages from terminal to neutral are in phase with each other; that is, in a three-phase Y connected generator, a single-phase voltage of triple frequency exists between the neutral and all three terminals, and the neutral therefore is not a true neutral. Between the lines no triple frequency voltage exists, since from terminal to neutral and from neutral to the other terminal the two third harmonics are in opposition and so neutralize.

This third harmonic between generator neutral and line must be kept in mind, since when large it may produce dangerous voltages by resonance with the line capacity.

When the generator neutral is grounded, the potential difference from line to ground is not line voltage divided by $\sqrt{3}$, that is, the true Y voltage of the system; but superimposed upon it is this single-phase triple frequency voltage; and the voltage from line to ground, especially its maximum, may be greatly increased, thus increasing the insulation strain. For this single-phase voltage all three lines go together, and so may cause static induction on other circuits, as telephone lines. A circuit of this single-phase triple frequency voltage then exists from the generator neutral over the inductance of all three generator circuits in multiple, and over the capacity of all three lines to ground, back to the generator neutral; that is, we have capacity and inductance in series in a circuit of the triple harmonic, and if capacity and inductance are high enough, we may get a dangerous voltage rise.

In this case of grounded generator neutral, if the neutral of the Y connected step-down transformers is grounded also, and the low tension side of these transformers connected in Y, the third harmonic of the generator has no path; the current produced by it would have to return over the open circuit reactance of the step-down transformer, and is limited thereby to a negligible value.

If, however, the secondaries of the step-down transformers are connected in delta, so that the third harmonic can circulate in the secondary delta, the third harmonic can flow through the transformer primary by inducing an opposite current in the secondary; in this case the step-down transformer short circuits the third harmonic of the generator. Grounding the primary neutral of step-down transformers

with grounded generator neutral therefore is permissible only if the transformer secondaries are also connected in Y. With delta connected transformer secondaries, however, it is not safe to ground the generator neutral and transformer neutral; since this produces a triple frequency current in generator, line and transformer; and even if the generator reactance is so high that the generator is not harmed by this current, it may burn out at the transformer, and probably will do so if the transformer is small compared with the generator.

This therefore is a case where delta connection of the transformer secondaries does not eliminate the trouble from the third harmonic, but makes it worse.

The triple frequency voltage from line to ground would be eliminated by short circuiting it in this manner, by Y delta connection of step-down transformer with grounded generator and transformer neutral, and static induction on other circuits so would disappear; but we get magnetic induction from the three triple frequency single-phase currents which now flow over the lines to the ground.

If the generator neutral is not grounded, it is safe to ground transformer neutrals. With ungrounded generator neutral, a triple frequency voltage can be measured by voltmeter, which then appears between generator neutral and ground; this voltage under unfavorable conditions, may give insulation strains in the generator by resonance rise; in the circuit from generator neutral over triple frequency voltage, generator inductance, capacity from line to ground and capacity from ground to generator winding in series.

In this case the capacity is much lower and the power therefore much less, that is, less danger exists.

When running two or more three-phase generators in parallel, with grounded neutrals:

a. If the generators have different third harmonics, these harmonics are short circuited from neutral over generator to the other generator and back to neutral; a triple frequency current thus flows between the generators, that is, the current between the generators can never be made to disappear.

That is, for the third harmonic, the two generators are two single-phase machines of different voltage, having the neutral as one terminal and the three three-phase terminals as the other single-phase terminal.

b. With two identical generators running in multiple, if the excitation is identically the same, no current flows between the grounded neutrals. If the excitation of the two generators is different, one is over-excited the other is under-excited (that is, one carries leading, the other lagging current) then a triple frequency current flows between the neutrals of identical generators. Since in parallel operation the terminal voltages are in phase, if by difference of excitation the two terminal voltages have a different lag behind the induced e. m. f.'s, the third harmonics, which lag three times as much as the fundamentals, cannot be in phase in the two machines; and thus triple frequency current flows between the machines.

In machines of very low reactance as turbo-alternators, even small differences in excitation of identical machines with grounded neutral may thus cause very large neutral currents.

In parallel operation of three-phase machines with grounded neutral, machines of different wave shapes frequently cannot be run together at all without excessive neutral currents, and the ground has to be taken off of one of the machine types.

Even with identical machines, such care has to be taken in keeping the same excitation that it is frequently undesirable to ground all the neutrals, but only the neutral of one machine is grounded and the other machine neutrals are left isolated. In

this case, provisions must be made to ground the neutral of some other machine, if the first one is out of service. The best way is, when grounding generator neutrals, to ground through a separate resistance for every generator and to choose this resistance so high as to limit the neutral current, but still low enough so that in case of a ground on one phase, enough current flows over the neutral to open the circuit breaker of the grounded phase.

The use of a resistance in the generator neutral is very desirable also, since it eliminates the danger of a high frequency oscillation between line and ground through the generator reactance in the path of the third harmonic, by damping the oscillation in the resistance. For this reason, the resistance should be non-inductive. To ground the generator neutral through a reactance is very dangerous since it intensifies the danger of a resonance voltage rise.

In grounding the generator neutral, special care is necessary to get perfect contact, since an arc or loose contact would generate a high frequency in the circuit of the third harmonic and so may lead to a higher frequency oscillation between line and ground.

SEVENTH LECTURE

HIGH FREQUENCY OSCILLATIONS AND SURGES

IN an electric circuit, in addition to the power consumption by the resistance of the lines, an energy storage occurs as electrostatic energy, or electrostatic charge due to the voltage on the line (capacity); and as electromagnetic energy, or magnetic field of the current in the line (inductance). In the long distance transmission line, both amounts of stored energy are very considerable, and of about equal magnitude; the former varying with the voltage, the latter with the current in the line. Any change of the voltage on the line, or the current in the line, or the relation between voltage and current, therefore requires a corresponding change of the stored energy; that is, a readjustment of the stored energy in the system, the electrostatic energy $\frac{e^2C}{2}$ and the electromagnetic energy $\frac{i^2L}{2}$, from the previous to the changed circuit conditions. This readjustment occurs by an oscillation, that is, a series of waves of voltage and of current, which gradually decreases in intensity, that is, dies out.

These oscillating voltages and currents are the result of the readjustment of the stored energy of the circuit to a sudden change of conditions, and are dependant upon the stored energy of the circuit, but not upon the generator frequency or wave shape; therefore they occur in the same manner, and are of the same frequency, in a 25 cycle system as in a 60 cycle system, or a high potential direct current transmission; and occur with sine waves of generator voltage equally as with distorted

generator waves. While the power of these oscillations ultimately comes from the generators, it is not the generator wave nor one of its harmonics which builds up, as discussed in the previous lectures; but the generator merely supplies the energy, which is stored as electrostatic charge of the capacity and as magnetic field of the inductance, and the readjustment of this stored energy to the change of circuit conditions then gives the oscillation.

These oscillating voltages and currents, adding to the generator voltage and current, thus increase the voltage and the current the more, the greater the intensity of the oscillation, and so may lead to destructive voltages.

Obviously, the intensity of the oscillation, that is, its voltage and current, are the greater, the greater or more abrupt the change was in the circuit, which caused the oscillation by requiring a readjustment of the energy storage. The greatest change in a circuit, however, is the change from short circuit to open circuit, and the instantaneous opening of a short circuit on a transmission line—as it occasionally occurs by the sudden rupture of a short circuiting arc—therefore gives rise to the most powerful, and thereby most destructive oscillation.

The wave length of oscillation thus depends on the length of the circuit in which the stored energy readjusts itself. For instance, in the short circuit oscillation of the system, the wave extends over the entire circuit, including generators and transformers; and the entire circuit so represents one wave, or one-half wave, that is, the wave length is very considerable. If the readjustment of stored energy takes place only over a section of the circuit, the wave length is shorter. For instance, if by a thunder cloud a static charge is induced on the transmission line, and by a lightning flash in the cloud, the cloud discharges, the electrostatic charge induced by it on the line

is set free and dissipates by an oscillation. In this case, the length of section on which an abnormal charge existed—one mile for instance—is a half wave of the oscillation, and the complete wave length would thus be two miles. Or, if a momentary discharge occurs over a lightning arrester to ground, the wave length may be only a few feet.

The velocity with which the electric wave travels in an overhead line is practically the velocity of light, or about 188,000 miles per second: it would be exactly the velocity of light, except that by the resistance of the line conductor the velocity is very slightly reduced. In an underground cable, by the high capacity of the cable insulation, the velocity of wave travel is greatly reduced, to about 50 to 70% of that of light.

From the wave length and the velocity follows the duration or time of one wave, and thereby the frequency of the oscillation. For instance, in the wave of two miles' length resulting from induction by a thunder cloud, as discussed above, the duration of the wave, or the time it takes to travel the wave length of two miles, at 188,000 miles per second velocity, is $\frac{2}{188,000} = \frac{1}{94,000}$ second, and thus, during one second, 94,000 waves would pass, that is, the frequency is 94,000 cycles. Or, if a transmission line of 80 miles' length short circuits at one end, and then disconnects at the other end by the opening of the circuit breaker, in the oscillation produced thereby the circuit is one-half wave. As the length of the circuit is $2 \times 80 = 160$ miles—conductor and return conductor,—the half wave is 160 miles; the complete wave therefore is $2 \times 160 = 320$ miles long, and the duration of the wave is $\frac{320}{188,000} = \frac{1}{587}$ seconds; the frequency 587 cycles, and if this

short circuit oscillation extends into, and includes the generating system, the frequency may be still lower.

Again, an oscillation of a very short section of the line, as for instance, 100 feet = $\frac{.00}{5280} = \frac{1}{52.8}$ miles wave length, would have a duration of the wave of $\frac{1}{52.8 \times 188,000} = \frac{1}{9,900,000}$ second, or a frequency of 9.9 millions of cycles per second.

Hence the frequency of such oscillations, caused by the readjustment of the stored energy of the system, may vary from values as low as machine frequency, up to many millions of cycles per second. It is the higher, the shorter the section of the circuit is in which the readjustment of energy occurs. The higher the frequency, and therefore the shorter the section of the circuit in which energy readjustment occurs, obviously the less is the amount of energy which is available in the oscillation—the stored energy of this section—and the less destructive therefore is the oscillation. That is, very high frequency oscillations are of very low energy and therefore of little destructiveness; but the energy and thus the destructiveness of an oscillation increases with decreasing frequency, and consequent increasing extent of the oscillation.

Such oscillations in a transmission line may result:

a. From outside sources, atmospheric electric disturbances, as illustrated in the above instance.

b. They occur during normal operation of the system: any change of load, or switching operation, as connecting or disconnecting circuits, etc., results in an oscillation, which usually is so small as to be harmless.

c. It may result from a defect or fault in the circuit, as an arcing ground or spark discharge, etc.

One of the most serious and destructive oscillations or surges is that produced by a spark discharge to ground, or an arcing ground, in an overhead transmission line or an underground cable system.

Assuming for instance a 44,000 volt transmission line of 50 miles' length, which is insulated from ground, that is, in which the neutral is not grounded. At 44,000 volts between the line conductors, the voltage between each conductor and the ground, normally, that is, with all conductors insulated, is $\frac{44,000}{\sqrt{3}} = 25,000$. If now somewhere in the middle of this line an insulator breaks, and the conductor thus drops near the grounded insulator pin or cross arm to about 2"; with 25,000 volts between conductor and ground, a spark would jump from the conductor to the ground, at the broken insulator, over the 2" gap. This spark develops into an arc, over which the electrostatic charge of the conductor discharges to ground as current, and the voltage of this conductor against ground thus falls to zero, since it is grounded by the arc; the two other line conductors then have the full line voltage, of 44,000, against ground; and their electrostatic charge against ground therefore increases, from that corresponding to their normal potential of 25,000, to that corresponding to 44,000 volts. As soon as the first conductor has discharged and fallen to ground potential, the current from this conductor to ground, over the gap, ceases, the arc goes out, and the conductor so is again disconnected from ground. It then begins to charge again to its normal potential of 25,000 volts against ground, while the other two conductors discharge, from 44,000 down to 25,000

volts. As soon, however, as during the charge the voltage of the first conductor has risen to the voltage required to jump across a 2" gap, this conductor again discharges to ground by a spark, which develops into an arc and so on, the phenomena of discharge and charge of the conductor repeating continuously. Such an oscillation, which continues indefinitely, that is, until the defect in the circuit is remedied, or the circuit has broken down and gone out of service, is usually called a *surge*. The duration of each oscillation of such an arcing ground is the time required: 1. To develop the arc, 2. to discharge the line, 3. to extinguish the arc, 4. to charge the line. In the above instance, the time of charge or discharge of the 25 miles of line from the arcing ground to the terminal station is: $\frac{25}{188,000} = \frac{1}{7520}$ second. Assuming the velocity of the arc stream as about 2000 feet per second, the development or extinction of a 2" arc would require $\frac{2}{12 \times 2000} = \frac{1}{12,000}$ second, and the total duration of one oscillation therefore is: $\frac{1}{12,000} + \frac{1}{7520} + \frac{1}{12,000} + \frac{1}{7520} = \frac{1}{2300}$ second, so giving a frequency of 2300 cycles.

The two other lines therefore oscillate in voltage against ground, that is, charge and discharge also at a frequency of 2300 cycles. They receive their charge, however, over the transformers at the two ends of the line, and their capacity therefore is in series with the self-inductance of these transformers in the circuit of the surge frequency of 2300 cycles; and the voltage of the other two lines thus may build up by the combination of capacity and inductance in series, to excessive values; that is, a destructive breakdown occurs from the other lines to ground—or in the apparatus connected to them in the terminal stations of the line, as transformers, current transformers, etc.

A spark discharge or oscillating ground therefore is one of the most serious, as well as not infrequent disturbances on a long distance transmission line or underground cable circuit; and it is mainly as a protection against this surge that it is recommended by many transmission engineers to ground the neutral of the system and so immediately convert a spark discharge on one conductor into a short circuit of one phase of the system, and thereby automatically cut out the circuit; that is, rather shut down this circuit than continue operation with an arcing ground on the system. Where, as in underground cable systems, a number of cables are used in multiple, the immediate disconnection of an arcing cable undoubtedly is advisable. In a single overhead transmission line, where a shutdown means a discontinuity of service, the question, whether by grounding the neutral it is preferable to shut down immediately in the case of an arcing ground, or continue service with ungrounded neutral and try to find and eliminate the arcing ground on the conductor, depends upon the length of time which the surge would probably last before causing a break down; it thereby depends upon the character of the circuit; the margin of insulation in transformers and insulators; and also on the value of continuity of service. The question of grounding or not grounding the neutral of a transmission line therefore requires investigation in each individual instance.

EIGHTH LECTURE

GENERATION

For driving electric generators the following methods are available :

1. The hydraulic turbine in a water power station.
2. The steam engine.
3. The steam turbine.
4. The gas engine.

COMPARISON OF PRIME MOVERS

1. The advantages of *water power, compared with steam power, are:*

a. Very low cost of operation : no fuel, very little attendance.

The disadvantages are :

a. Usually the cost of development and installation is far higher than with steam power.

b. The location of the water power cannot be chosen freely, but is fixed by nature; therefore the power cannot be used where generated, but a long distance transmission line is required.

c. Usually lower reliability of service, due to the dependence on a transmission line, and on meteorological conditions: the river may run dry in summer, ice interfere with the operation in winter.

The speed of the water in the turbine depends upon the head of water, and is approximately, in feet per minute, $480\sqrt{h}$, where h is the head, in feet. The peripheral speed of the turbine, and so its revolutions, depends upon the speed and therefore upon the head of the water. At high heads of 500 to

2000 feet, as are found in the West, the electric generators are thus high speed machines, of good economy and moderate size and cost. At low heads, however, such as are usual in the Eastern States, direct connection to a turbine leads to slow speed generators of many poles and large size and cost; while indirect driving, by belt or rope, is mechanically undesirable. Very low head water powers of less than 20 to 30 feet head therefore are of little value and their development is economical only where electric power is valuable.

Of the two types of turbines, the reaction turbine runs approximately at the speed of the water, and the action or impulse turbine at half the speed of the water. At the same head and thus the same speed of the water, the reaction turbine gives higher speed, and is therefore used in water powers of low and medium heads, where the speed of the water is low; while the impulse turbine, as the Pelton wheel, is always used at very high heads, at which the reaction turbine would give too high speeds.

Where water power is not available, the power has to be generated by the combustion of fuel. In this case, a greater freedom exists in the choice of the location of the plant; and it is located as near to the place of consumption as considerations of the cost of property, the availability of condensing water for the engines, the facilities of transportation, etc., permit. Transmission lines therefore are less frequently used, but in steam stations of large power, high potential distribution circuits of 6600, 11,000 or 13,200 volts, commonly underground by cables, are used in supplying electric power from the main generating station, to the substations as centres of secondary distribution (New York, Chicago, etc.).

As source of power is available then:

The steam engine. The steam turbine. The gas engine.

Comparison of the steam turbine with the steam engine:

Some of the advantages of the steam turbine over the steam engine are:

- a. High efficiency at low loads, and a flatter efficiency curve; that is, the turbine efficiency remains high at partial loads, and at overloads, where the steam engine efficiency falls off greatly; so that the superiority of the steam turbine in efficiency, while marked at rated load, is still far greater at partial load, light load and overload.
- b. Smaller size, weight and space occupied.
- c. Uniform rate of rotation, therefore decreased liability of hunting of synchronous machines, and decreased necessity of heavy foundations to withstand reciprocating strains.
- d. Greater reliability of operation and far less attendance required.

The steam turbine reaps a far greater benefit in economy than the steam engine from superheat of the steam, and from a high vacuum in the condenser.

Some of the disadvantages of the steam turbine are:

- a. It is a new type of machine, developed only within the last ten years, and operating engineers and attendants are therefore less familiar with it than with the reciprocating engine; and the steam turbine is replacing the steam engine in electric power plants so rapidly, that it is difficult to get sufficient men to intelligently install and operate them.

It is therefore of greatest benefit in a steam turbine installation that the user familiarize himself with the machine, so as not to depend upon the manufacturer in every minute detail, but take care of minor troubles just as he would do with a steam engine. As the steam turbine is a very simple apparatus this is not difficult.

The speed characteristic of the steam turbine is similar to that of the constant voltage direct current shunt motor, or the polyphase induction motor; while that of the reciprocating steam engine is similar to that of the series motor. That is, to produce the same torque, the steam turbine requires approximately the same amount of steam, irrespective of the speed; therefore its efficiency is highest at a certain speed, or rather range of speed, but falls off with the speed; while the steam consumption of the reciprocating engine, at constant torque, is approximately proportional to the speed, that is the number of times the cylinders are filled per minute. Or in other words, the torque per pound of steam used per minute is approximately constant and independent of the speed in the turbine (just as the torque per volt-ampere is approximately constant for all speeds in the induction motor), while in the reciprocating engine the torque per pound of steam used per minute is approximately inversely proportional to the speed, or at least greatly increases with decrease of speed (just as in the series motor the torque per volt-ampere input increases with decrease of speed).

The steam turbine therefore would not be suitable for directly driving a railway train in rapid transit service, but is suitable for driving the ship's propeller.

Just as in the induction motor a series of economical speeds can be produced by changing the number of poles, so in the steam turbine a series of economical speeds can be produced by changing the number of expansions. For driving electrical machinery this, however, is of no importance.

Comparison of the gas engine with the steam turbine and the steam engine.

The leading and foremost advantage of the gas engine, and the feature which gives it the right of existence, is its

high efficiency. That is, the same amount of coal, converted to gas and fed to a good gas engine, gives far more power than when burned under the boilers of the most efficient steam turbine. The cause is that the gas engine works over a far greater temperature range than the steam engine and even the steam turbine—although the latter, by its ability to economically utilize superheat and high condenser vacuum, gets the benefit of a larger temperature range over the steam engine.

If therefore the gas engine were not so very greatly handicapped in every other respect, it would long have superseded the steam engine and the steam turbine.

The disadvantages of the gas engine in every respect but efficiency are such, however, that in spite of its existence of over half a century it has not made a serious impression on the industry; while the steam turbine in the last ten years of its development has practically replaced the steam engine in large electric generating plants.

The cause of the disadvantages of the gas engine is the high maximum temperature and the high maximum pressure compared with the mean pressure in the cylinders, which is necessary to get the greater temperature range and thus the efficiency, therefore is inherent in this type of apparatus.

The output depends upon the mean pressure in the cylinder, which is low; the strains on the maximum pressure, which is very high; and the gas engine therefore must be very large, and its moving parts very strong and heavy, for its output. The impulse due to the rapid pressure change is very jerky—almost of the nature of an explosion—and the steadiness of the rate of rotation is therefore very low, requiring for electric driving very heavy flywheels and numerous cylinders.

Compared with the steam engine, the disadvantages of the gas engine so are:

- a. Lower reliability; higher cost of maintenance in attendance, repairs, and greater depreciation.
- b. Larger size and space occupation for the same output.
- c. Less ease to start.
- d. In general, lower steadiness of the rate of rotation.

The advantage of the gas engine is, that it requires no boiler plant; the compensating disadvantage, that it requires a gas generating plant. This latter disadvantage disappears where gas is available as fuel—in the waste gases of blast furnaces of steel plants and in the natural gas districts—and in those cases gas engines have found their introduction. They have also been installed for smaller powers, where low cost of fuel is unessential, but the operation of a steam boiler is objectionable, as in isolated plants using city gas or liquid fuel (gasolene, etc.).

In general, however, with the exception of those special cases, the gas engine does not yet come into consideration in the electric power generating station.

ELECTRIC GENERATORS

In general, considerations of economy make it desirable to generate the electric power in the form in which it is used. In most cases, however, this is not feasible, but a higher voltage or even a different form of power (alternating instead of direct) is necessary in the generating station than that required by the user, to enable transmission and distribution; and then usually three-phase alternating current is generated.

1. For isolated plants, and in general distribution of such small extent as to be within range of 220 volt distribution, 220 volt direct current generators are used, operating a three-wire system, either two 110 volt machines, supplying the two sides of the system, or 220 volt machines, deriving the

neutral by equalizer machines, or by connection to a storage battery, or by compensator and collector rings on the 220 volt generator. That is, two diametrically opposite (electrically) points of the armature winding are connected to collector rings, (so giving an alternating current voltage on those collector rings), an alternating current compensator (transformer with a single winding) is connected between the collector rings, and the neutral brought out from the center of the compensator, as shown diagrammatically in Fig. 24. This arrangement is now most commonly used.

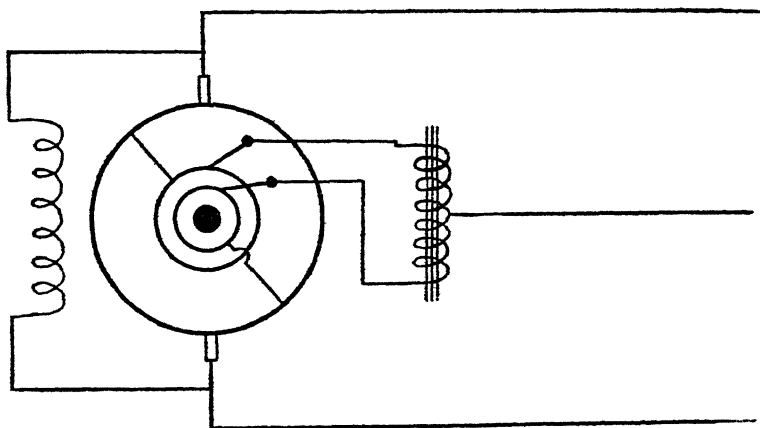


Fig. 24

For direct current distribution in larger cities, such generating stations have practically disappeared, and have been replaced by converter substations, receiving power from a 6600, 11,000 or 13,200 volts, and usually 25 cycles.

2. For street railway, 600 volt direct current generators main generating station, as three-phase alternating current of are still used to a considerable extent, where the railway system is of moderate extent. In large railway systems, and roads covering greater distances, as interurban trolley lines,

direct generation of 600 volts direct current is also disappearing before the railway converter substation, receiving power as three-phase alternating from transmission lines or high voltage distribution cables.

3. For general distribution by alternating current, with a 2200 volt primary system, direct generation is still largely used, as the use of 2200 volt permits the system to cover a very large territory, and substations are mainly used only where the power can be derived from a long distance transmission line, or where the 2200 volt distribution is only a part of a large system of electric generation; as in the suburban distribution of large cities, using converter substations for the interior. In this case, where the transmission line or the main generating station is at 60 cycles, large station transformers are used for the supply of the 2200 volt distribution; where the power supply is at 25 cycles, either frequency converters, or motor generators change to 60 cycles, 2200 volts.

4. For special use, as for electrochemical work, where the electric power is generated directly, different voltages, etc., may be used to suit the requirements.

Where the power cannot be generated in the form in which it is used, and that is the case in all larger systems, three-phase alternators are almost universally used.

The single-phase system has the disadvantage that single-phase induction and synchronous motors and converters are inferior to polyphase machines, and single-phase alternators larger and less efficient, and for lighting, where single-phase is preferable, single-phase lighting circuits can be operated from polyphase alternators.

Two-phase also is gradually going out of use, since it offers no advantage over the three-phase, and the three-phase is

preferable for transmission, requiring only three conductors, while two-phase requires four.

In polyphase alternators the flow of power is constant, that is, at any moment adding the power of all phases gives the same value, while in single-phase alternators the power is pulsating.

In a polyphase machine the armature reaction also is constant, in a single-phase machine, pulsating; in the latter therefore, in machines of very large armature reaction, as turbo-alternators, pulsations of the magnet field, and thereby loss in efficiency, and heating may result.

An alternator has armature reaction and self-induction.

The armature reaction is the magnetic action of the armature current on the field, that is, the armature current demagnetizes or magnetizes the field according to its phase, and so lowers or raises the voltage. Armature reaction therefore is expressed in ampere turns.

Self-induction is the action of the armature current in producing magnetism in the armature, which magnetism does not go through the field. This magnetism induces an e. m. f. in the armature, which opposes or assists the e. m. f. produced by the field magnetism, according to the phase of the armature current, and so lowers or raises the voltage. Self-induction, or "*armature reactance*" therefore is expressed in ohms.

Armature reaction and self-induction therefore act in the same manner, lowering the voltage with lagging and raising the voltage with leading current.

In calculating alternators, either the armature reaction and the self-induction can both be considered, which makes the calculation more complicated; or the armature reaction may be neglected and the self-induction made so much larger as to allow for the armature reaction. This self-induction is then

called the "synchronous reactance" and, combined with the armature resistance, the "synchronous impedance" of the machine. Or the self-induction may be neglected and only the armature reaction considered, but which is increased to allow for the self-induction.

The last way (armature reaction), is used in designing machines; the second way (synchronous reactance) in calculations with machines and systems.

In the momentary short circuit current of alternators, however, the armature reaction and the self-induction must be considered separately, since they act differently.

In the moment of short circuiting an alternator, the self-induction acts immediately in limiting the current, but not so the armature reaction, because it takes time before the armature current demagnetizes the field, that is, the field exciting winding acts as a short circuited secondary around the field poles, and retards the decrease of field magnetism resulting from the demagnetizing action of the armature current by inducing a current in the field winding, which tends to maintain the field magnetism.

Therefore in the first moment after the short circuit the armature current is limited by self-induction only, and is therefore much larger than afterwards, when self-induction and armature reaction both act.

In machines of low armature reaction and high self-induction, as high frequency alternators, the momentary short circuit current is not much larger than the permanent short circuit current. In machines of low self-induction, that is, of a well distributed armature winding, but high armature reaction, (that is, very large output per pole, as in steam turbine alternators,) the momentary short circuit current may be many

times greater than the permanent value of the short circuit current, which is reached after a few seconds.

In the moment of short circuiting such an alternator, the field current rises to several times its normal value, and becomes pulsating, of double frequency. Gradually the armature current and the field current die down to their normal values. By inserting non-inductive resistance in the field circuit of the alternator, the field current, which is induced in the moment of short circuit, can be forced to die out more rapidly, and the armature short circuit current made thereby to reach its final value more quickly, that is, the duration of the excessive momentary short circuit current may be reduced.

By inserting reactance, as choke coils or reactive coils, in the armature circuit of the alternator, its momentary short circuit current can be reduced, and this is advisable in such machines in which the current otherwise would reach dangerous values. Since the regulation of such alternators mainly depends upon the armature reaction, which is very large compared with the self-induction, even a considerable external self-induction inserted as reactive coil for limiting the momentary short circuit current does not much increase the combined effect of armature reaction and self-induction; that is, does not seriously affect the regulation.

NINTH LECTURE

HUNTING OF SYNCHRONOUS MACHINES

CROSS currents can flow between alternators due to differences in voltage, that is, differences in excitation; and due to differences in phase, that is, differences in position of their rotors.

Cross currents due to differences in excitation are wattless currents, magnetizing the under-excited and demagnetizing the over-excited machine.

Cross currents due to differences in position are energy currents, accelerating the lagging and retarding the leading machine. Their magnetic action is a distortion or a shift of the field, that is, they increase the magnetic density at the one and decrease it at the other pole corner.

If two machines are thrown together out of phase, or brought out of the phase by some cause (as the beat of an engine, or the change of load of a synchronous motor) then the two machines pull each other in phase again, oscillate a few times against each other, which oscillation gradually decreases and dies out, and the machines run steadily.

If the oscillations do not decrease, but continue, the machines are said to be hunting.

If the oscillation is small it may do no harm; if it is greater, it may cause fluctuation of voltage, resulting in flickering of lights, etc.; if it gets very large, it may throw the machines out of step.

Some causes of hunting are:

- 1st. Magnetic lag.
- 2nd. Pulsation of engine speed.
- 3rd. Hunting of engine governors.
- 4th. Wrong speed characteristic of engine.

1st. When the machines move apart from each other, magnetic attraction opposes their separation. When they pull together again, magnetic attraction pushes them together with the same force, so that they would move over the position of coincidence in phase and separate again in the opposite direction just as much as before.

Energy losses as friction, etc., retard the separation and so make them separate less than before, every time they do so, that is, cause them gradually to stop see-sawing.

If, however, there is a lag in the magnetic attraction, then they come together with greater force than they separated, so separate more in the opposite direction, that is, the oscillation increases until the machines fall out of step, or the further increase of oscillation is stopped by the increasing energy losses.

This kind of hunting is stopped by increasing the energy losses due to the oscillation, by copper bridges between the poles, by aluminum collars around the pole faces, or by a complete squirrel cage winding in the pole faces.

The frequency of this hunting depends on the magnetic attraction, that is, on the field excitation, and on the weight of the rotating mass. The higher the field excitation the greater is the magnetic force, that is, quicker the motion of the machine and therefore the higher the frequency. The greater the weight, the slower it is set in motion, that is, the lower the frequency.

Characteristic of this hunting therefore is that its frequency is changed by changing the field excitation.

2nd. If the speed of the engine varies during the rotation, rising and falling with the steam impulses, then the alternator speed and the frequency also pulsate with a speed equal to, or a multiple of the engine speed. If now two

such alternators happen to be thrown together so that the moment of maximum frequency of one coincides with the moment of minimum frequency of the other, the two machines cannot run in perfect phase with each other, but pulsate, alternately getting out of phase with each other, coming together, and getting out again in the opposite direction. If the deviation of the two engines from uniform rate of rotation is very little—the maximum displacement of the alternator from the position of uniform rotation not more than three electrical degrees—the pulsating cross currents, which flow between the alternators, are moderate, and the phenomenon harmless, as long as the oscillation is not cumulative. An increase of the weight of the flywheel of the engine decreases the speed pulsation and thereby decreases this form of hunting, which is the most harmless, but increases the tendency to the hunting in No. 1 and No. 3, and therefore is not desirable; but steadiness of engine speed should be secured by the design of the engine, that is, by balancing the different forces in the engine, as the steam impulses and the momentum of the reciprocating masses, so as to give a uniform resultant.

In such a case, when running from a single alternator, driven by a reciprocating engine with moderate speed pulsation, (therefore receiving a slightly pulsating frequency) a synchronous motor without anti-hunting devices, but of high armature reaction, and therefore high stability, may run very steadily, with no appreciable current pulsation; while the same synchronous motor, when supplied with a squirrel cage winding in the field pole faces as the most powerful anti-hunting device, may show pulsation in the current supplied, which in a high speed motor, of high momentum, may be considerable. The cause is, that in the former case the synchronous motor does not follow the pulsation of frequency, but keeps constant

speed, while in the latter case the squirrel cage winding forces the motor to follow the variation in frequency by accelerating and decelerating, and the pulsation of the current therefore is not hunting, but energy current required to make the motor speed follow the engine pulsation.

If the frequency of oscillation of the machine (as determined by its field excitation and the weight of its moving part) is the same as the frequency of engine impulses, that is, the same as the number of engine revolutions or a multiple thereof, then successive engine impulses will always come at the same moment of the machine beat and so continuously increase it: that is, the machine oscillation increases, or the machine hunts.

In this case of cumulative hunting caused by the engine impulses, the frequency of oscillation agrees with the engine oscillation.

3rd. If one alternator is a little ahead, that is, takes a little more load, its engine governor regulates by reducing the steam, slowing down the alternator to its normal position. When slowing down, the flywheel is giving power, therefore the steam supply has been reduced more than it should be, that is, the alternator drops behind and takes less load until the governor has admitted steam again.

In the meantime, while the first alternator was behind and took less load, the second alternator had to take the load, that is, the governor of the second alternator admitted more steam. When the first alternator has picked up again to its normal load, the second alternator gets too much steam and its governor must cut off, but then cuts off too much, the same way as the first alternator did before; so the two governors hunt against each other by alternately admitting too much and too little steam.

In this case the frequency of hunting does not depend on the engine speed and does not vary much with the field excitation, but the hunting is usually much less at heavy load than at light load. The reason is that at load, when the engines take much steam, a little change in the steam supply does not make so much difference as at light load, where the engines take very little steam, and so a small change of the governor has a great effect.

4th. To run in parallel, the speed of the engines driving the alternators must decrease with the load so that the alternators divide the load.

If the speed did not change with the load, then there would be no division of the load; the one engine could take all the load, the other nothing.

If the speed curve of the engine is such that the speed does not fall off much between no load and moderate load, then the alternators will not well divide the load at light loads, and hunt while running in parallel at light load, but steady down at heavier loads.

To distinguish between different kinds of hunting:

1st. Change of frequency with change of field excitation points to magnetic hunting, especially if very marked.

2nd. Equality of frequency with the generator speed points to engine hunting.

3rd. If the synchronous motor or converter steadies down when only one engine is running, it points to engine governor hunting.

4th. Steadiness of operation at load, and unsteadiness at light load points to governor hunting, but may also be due to engine and magnetic hunting.

5th. If by disconnecting one governor and governing one engine only, the hunting disappears, then it is due to

governor hunting. If it does not disappear, then both governors may be disconnected and the engines run carefully without governors, by throttle. If the hunting then disappears, it is due to the governors; if it does not disappear, it is probably magnetic hunting.

If by making the field excitation of the two alternators or two converters that hunt, unequal—by increasing the one and decreasing the other—the hunting disappears or decreases, it is magnetic hunting.

In a case of hunting, the following points should be investigated:

A. HUNTING OF SYNCHRONOUS MOTORS OR CONVERTERS

1st. Count the number of beats to get the frequency of hunting. If the beats periodically increase and decrease, it shows two frequencies of hunting superimposed upon each other. Then count the total number of beats per minute (counting during intermissions) and count the number of intermissions per minute.

The two frequencies are the number of beats per minute, plus and minus half the number of intermissions or nodes per minute.

Instance: 80 beats per minute, 10 intermissions per minute. Frequencies $80 + 5$ and $80 - 5$ or 85 and 75 beats.

If one of the two frequencies approximately coincides with the engine speed, it can be assumed as the engine speed. The number of revolutions of the engine obviously should be counted also.

2nd. See whether any machine in the system runs at a speed equal to the observed frequency of hunting. For instance, a generator may make 75 revolutions per minute, which accounts for this frequency.

3rd. With several converters in the same station see whether the station ammeter also hunts.

If the station ammeter is very steady and the converter ammeters hunt, the converters hunt against each other. In this case lowering the one and raising the other converter field and, if necessary, readjusting the potential regulators, may stop the hunting by giving the two machines different frequencies of hunting which interfere with each other.

If all three meters are unsteady, the converters may hunt against each other or hunt together against another station or against the generator. Then find out whether the ammeter needles of both converters go up and down together or one goes up when the other goes down.

4th. Change the field excitation and see whether the change of field excitation changes the frequency. See whether a decrease of field excitation steadies it. Occasionally hunting can be stopped by lowering the field excitation, that is, running with lagging current.

5th. If several converters of a substation feed into the same direct current system, as the converters of other substations, disconnect the direct current sides of the converters and see if they still hunt.

If two or more converters run in the same station, run only one and see whether it hunts.

CURE

1st. If the hunting is magnetic hunting between converters or synchronous motors, it is frequently reduced by making the field excitation unequal, or putting a flywheel on one converter, or belting some other machine to it, or running an induction motor in the same station or in any other way breaking up the resonance.

2nd. Several converters hunting against each other in the same substation are frequently steadied by connecting the collector rings with each other, that is, by equalizer connections between converter and transformer or regulator.

In this case the commutator brushes have to be carefully adjusted to avoid sparking.

3rd. The most effective way is to put copper bridges on the converters or synchronous motors, or better still a squirrel cage winding in the field pole faces.

Not so good are short circuiting rings around the field poles.

B. HUNTING OF GENERATORS

1st. Count the frequency in the same way as before.

2nd. See whether the frequency agrees with the generator speed or with the speed of some large motor on the system.

3rd. See whether the frequency changes with the excitation.

4th. See whether the hunting changes with the load, that is, gets worse at light load.

5th. Disconnect governors and see whether this stops hunting.

CURE

1st. If the hunting stops when disconnecting the governors, it is hunting of the governors and can be cured by putting a stiff dashpot on the governors.

2nd. If the hunting does not stop by disconnecting the governors, copper bridges on the alternators will cure it.

3rd. If the hunting has the speed of the engine, it may be reduced by increasing the flywheel or decreasing it, by running an induction motor in the station, or in any other way breaking up the resonance.

In general, systems having all kinds of loads, different sizes of generators, motors and converters, induction motors and synchronous motors mixed, etc., are very little liable to hunting. Hunting is most liable to occur when all the generators are of the same kind and all the synchronous motors or converters are of the same kind.

Resistance between the machines increases the tendency to hunting so that if the resistance drop is more than 10% to 15%, special precautions have to be taken, such as squirrel cage pole face windings, or synchronous machines must be altogether avoided and induction motor generator sets used.

Reactance in general reduces the tendency to hunting except when very large.

The tendency to hunting is very severe at the end of a long distance transmission line and induction machines as a rule are preferable in such a place.

Machines with high armature reaction are much less liable to hunt than machines with low armature reaction, that is, close regulation, because with high armature reaction the current varies much less with a change of position of the machine. Therefore, 60 cycle converters are more liable to hunt than 25 cycle converters, because in 60 cycle converters there is not enough space on the armature to get high armature reaction.

TENTH LECTURE

REGULATION AND CONTROL

A. DIRECT CURRENT SYSTEMS.

In *direct current three-wire 220 volt distribution systems* several outside bus bars are used and, with change of load, the feeders are changed from one bus bar to another.

The different bus bars are connected to different machines, to the storage battery or to boosters.

The lighting boosters are low voltage machines separately excited from the bus bars. The main generators are shunt machines or rather are excited from the bus bars, or rotary converters, and are usually of 250 volts, that is, the neutral brought out by collector rings and compensator.

In railway circuits, in addition to trolley wire and rail return, trolley feeders and ground feeders, or plus and minus feeders are sufficient for converter substations, and where the distance gets too great for feeders, another substation is installed.

When using direct current generators, series boosters are used to feed very long feeders which otherwise would have an excessive drop of voltage. In this way feeder drops of 200 to 300 volts are taken care of by the railway booster. Such a large voltage drop is uneconomical and railway boosters are therefore used only for small sections for which it does not pay to install a separate station, especially where the load is very temporary, as for instance, heavy Sunday load, etc.

Railway boosters are series machines, that is, the series field and the machine voltage therefore are proportional to the current. In such railway boosters it is necessary to take care in the booster design that it does not build up as series generator

feeding a current through the local circuit between a short feeder and a long feeder, as shown in Fig. 25.

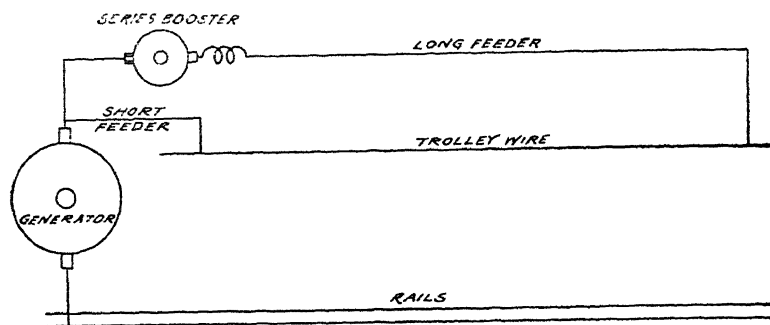


Fig. 25.

A series machine excites if the resistance of its circuit is less than a certain critical value. To avoid such local circuit, either the trolley circuit is cut between the feeders, or the boosting kept below the critical value.

If the distances are too great for boosters, inverted converters in the generating station are used to change from direct current to alternating current; the alternating current is sent by step-up and step-down transformers to the substation and changed to direct current by rotary converters.

If a considerable amount of power is required at a distance, it is more convenient at the generating station to use, instead of inverted converters, double current generators, that is, generators having commutator and collector rings.

If most of the power is used at a distance, alternating current generators are used with rotary converters and frequently one converter substation is located in the generating station.

Inverted converters and double current generators are now used less, since usually the systems are now so large as to

require most of the power at a distance, and therefore alternating current generators are used.

Many big systems have advanced from direct current generators, through inverted converters and double current generators, to the present alternators feeding converter substations.

B. ALTERNATING CURRENT SYSTEMS.

Generator Regulation.

1st. Close inherent regulation.

This is secured by low armature reaction and high saturation so that the voltage does not vary much with the load.

Advantages—

Simple, requiring no additional apparatus, etc.

Instantaneous.

Disadvantages—

Larger and more expensive generators and when of very close regulation, more difficult to run in parallel.

2nd. Rectifying Commutator.

The main current goes over a commutator, is rectified, and the rectified current sent through a series field. This arrangement is not used any more.

Advantage—

Permits compounding and over-compounding without any elaborate apparatus.

Disadvantages—

Only limited power can be rectified, therefore suitable only for smaller machines.

Compounds correctly only for constant power factor ; that is, if compounded for non-inductive load, the voltage drops on inductive load, since inductive load requires a greater field excitation than non-inductive load.

Brushes have to be shifted with change of power factor, that is, change from motor load to lighting load, etc.; otherwise commutator sparks badly.

These machines therefore were good in the early days when all the load was lighting load, but are unsuited at present for mixed load.

3rd. Form D alternator or compensated alternator with compensating exciter.

Exciter is connected direct and has the same number of poles as the alternator so as to be in synchronism.

The main current passes by collector rings through the exciter armature, usually with interposition of a transformer to keep the high voltage away from the exciter.

The main current is sent through the exciter in such direction that with non-inductive load its armature reaction slightly magnetizes and so raises the exciter voltage. Then with lagging current it magnetizes much more, raises the exciter voltage more, and with leading current demagnetizes or lowers the exciter voltage.

By adjusting the machine it can be made to compound perfectly for non-inductive, inductive, and leading current load and also can be made to over-compound at constant speed; so that with the engine dropping in speed by the load, it keeps constant voltage, or raises the voltage as desired.

Advantages—

Very quick and very correct compensation, and if properly adjusted, the most perfect arrangement and not liable to get out of adjustment.

Disadvantages—

More expensive machine and the average engineer not skillful enough to adjust it.

4th. Compensating Commutator (Alexanderson).

Commutator built multi-segmental so that it does not spark with fixed position of brushes.

Brushes set to rectify completely at inductive load; at non-inductive load then rectify incompletely, and so the series field gets less current.

At leading current load the rectified current is reversed and so the series field demagnetizes.

Advantages—

Takes care of lag and lead.

Disadvantages—

Special generator and suitable only for small and moderate sizes.

5th. POTENTIAL REGULATOR.

Tirrill Regulator

Rheostat in exciter field so large that when in circuit the excitation is the lowest, and that when short circuited the excitation is the highest ever required.

A potential magnet in the alternator circuit operates a contact maker which continuously cuts the resistance in and out again, so that the contact maker is never at rest, but always cuts in and out, and the average field excitation of the exciter is between maximum and minimum.

If the voltage tends to drop, the contact remains a shorter time on the low than on the high position, and so raises excitation; if the voltage tends to rise, the contact maker remains a shorter time on the high than on the low position, and so lowers excitation.

Advantages—

Very simple.

Can be applied to any alternator and requires no special adjustment.

Disadvantages—

Limited in power by the sparking of the contact maker, and so can be used only if the regulation of the alternator is not too bad or the machine too large.

In very large machines usually no regulating device is used but hand control of the field rheostat, since in such large machines the load only varies slowly and never changes much, as for reasons of economy the machines are run near full load; with the change of load, machines are shut down or started up.

Synchronous Motors and Converters.

In an alternating current system or part of the system containing large synchronous motors or converters the voltage can be controlled by varying the motor or converter field in the same way as with alternators, that is. by Tirrill Regulator or commutator and series field. etc.

POTENTIAL REGULATORS.

1. Compensator regulator.

With step-up or step-down transformers the voltage can be regulated by having different taps brought out of the transformer winding and so get different voltages by means of a dial switch. Where no transformers are used a compensator with different voltage taps gives the same results.

The taps can be brought out in the primary or in the secondary, whichever is the most convenient: in the secondary, if the primary is of very high voltage; in the primary, if the secondary is of very low voltage and large current.

Advantages—

Simplest, cheapest and most efficient.

Disadvantages—

Step by step variation and sparking at the dial switch.

2. INDUCTION REGULATORS.

Built like induction motors with stationary primary in shunt and movable secondary in series to the line.

By moving the secondary the voltage varies from lowering to raising.

Induction regulators are usually three-phase and of larger sizes for rotary converters in lighting systems.

When single-phase, the stationary member contains a short circuited coil at right angles to the primary. In the neutral position this coil acts as short circuited secondary to the secondary coil, and so reduces its self-induction.

Advantages—

Perfectly uniform variation and considerable inductance which is of advantage for rotary converters.

Disadvantages—

High cost.

3. MAGNETO REGULATORS.

A stationary primary coil is in shunt and a stationary secondary coil is in series and at right angles to the primary; an iron shuttle moves inside of the coils and so turns the magnetism of the primary coil into the secondary coil either one way or the other.

On the dotted position the primary sends the magnetism through the secondary in opposite direction as in the drawn position, in Fig. 26.

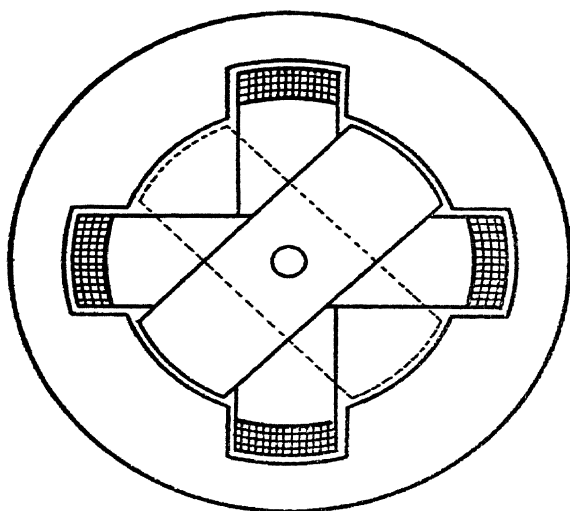


Fig. 26

Advantage—

Uniform variation.

Disadvantage—

More expensive than compensator regulator.

ELEVENTH LECTURE

LIGHTNING PROTECTION

WHEN the first telegraph circuits were strung across the country, lightning protection became necessary, and was given to these circuits at the station by connecting spark gaps between the circuit conductors and the ground.

When, however, electric light and power circuits made their appearance, this protection against lightning by a simple small spark gap to ground became insufficient, and this additional problem arose: to open the short circuit of the machine current, which resulted from and followed the lightning discharge.

This problem of opening the circuit after the discharge was solved by the magnetic blow-out, which is still used to a large extent on 500 volt railway circuits; by the horn gap arrester—a gap between two horn-shaped terminals, between which the arc rises, and so lengthens itself until it blows out; and later on, for alternating current, the multi-gap between non-arcing metal cylinders, a number of small spark gaps in series with each other, between line and ground, over which the lightning discharges to ground—the machine current following as arc, but stopped at the end of the half wave of alternating current; but not starting at the next half wave, due to the property of these “non-arcing” metals (usually zinc-copper alloys), to carry an arc in one direction, but requiring an extremely high voltage to start a reverse arc.

These lightning arresters operated satisfactorily with the smaller machines and circuits of limited power used in the earlier days, but when large machines of close regulation, and therefore of very large momentary overload capacity were in-

troduced, and a number of such machines operated in multiple, these lightning arresters became insufficient: the machine current following the lightning discharge frequently was so enormous that the circuit did not open at the end of the half wave, but the arrester held an arc and burned up.

Furthermore, the introduction of synchronous motors, and of parallel operation of generators, made it essential that the lightning arrester should open again instantly after discharge. For, if the short circuit current over the arrester lasted for any appreciable time: a few seconds, synchronous motors and converters dropped out of step, the generators broke their synchronism, and the system in this way would be shut down. The horn gap arrester, in which the arc rises between horn-shaped terminals, and by lengthening, blows itself out, therefore became unsuitable for general service; since without series resistance, the short circuiting arc lasted too long for synchronous apparatus to remain in step, and with series resistance reducing the current so as not to affect synchronous machines, it failed to protect under severe conditions. Thus it has been relegated for use as an emergency arrester on some overhead lines, to operate only when a shutdown is unavoidable.

To limit the machine current which followed the lightning discharge, and so enable the lightning arrester to open the discharge circuit, series resistance was introduced in the arrester. Series resistance, however, also limited the discharge current, and with very heavy discharges, such lightning arresters with series resistance failed to protect the circuits, that is, failed to discharge the abnormal voltage without destructive pressure rise. This difficulty was solved by the introduction of shunted resistances, that is, resistances shunting a part of the spark gaps. All the minor discharges then pass over the resistances and the unshunted spark gaps, the

resistance assisting in opening the machine circuit after the discharge. Very heavy discharges pass over all the spark gaps, as a path without resistance, but those spark gaps which are shunted by the resistance, open after the discharge; the machine current, after the first discharge, therefore is deflected over the resistances, limited thereby; and the circuit so finally opened by the unshunted spark gaps.

With the change in the character, size and power of electric circuits, the problem of their protection against lightning thus also changed and became far more serious and difficult. Other forms of lightning, which did not exist in the small electric circuits of early days, also made their appearance, and protection now is required not only against the damage threatened by atmospheric lightning, but also against "lightning" originating in the circuits: so called "internal lightning," which is frequently far more dangerous than the disturbances caused by thunder storms.

Under lightning in its broadest sense we now understand all the phenomena of electric power when beyond control.

Electric power, when getting beyond control may mean excessive currents, or excessive voltages. Excessive currents are rarely of serious moment: since the damage done by excessive currents is mainly due to heating, and even very excessive currents require an appreciable time before producing dangerous temperatures. Usually circuit breakers, automatic cut-outs, etc., can take care of excessive currents, and such currents produce damage only in those instances where they occur at the moment of opening or closing a switch, by burning contacts, or where the mechanical forces exerted by them are dangerously large, as with the short circuit currents of the modern huge turbo-generators.

Excessive voltage, however, is practically instantaneous in its action, and the problem of lightning protection therefore is essentially that of protecting against excessive voltages.

The performance of the lightning arrester on an electric circuit is analogous to that of the safety valve on the steam boiler, that is, to protect against dangerous pressures—whether steam pressure or electric pressure—by opening a discharge path as soon as the pressure approaches the danger limit. Therefore absolute reliability is required in its operation, and discharge with as little shock as possible, but over a path amply large to discharge practically unlimited power without dangerous pressure rise.

However, the causes of excessive pressures, and the forms which such pressures may assume, are so much more varied in electric circuits than with steam pressures, that the design of perfectly satisfactory lightning arresters has been a far more difficult problem than the design of the steam safety valve.

Such excessive pressures may enter the electric circuit from the outside by atmospheric disturbances as lightning, or may originate in the circuit.

Excessive pressures in electric circuits may be single peaks of pressure, or “strokes” or discharges, or multiple strokes; that is, several strokes following each other in rapid succession, with intervals from a small fraction of a second to a few seconds, or such excessive pressures may be practically continuous, the strokes following each other in rapid succession, thousands per second, sometimes for hours.

Atmospheric disturbances, as cloud lightning, usually give single strokes, but quite frequently multiple strokes, as has been shown by the oscillograms secured of such lightning discharges from transmission lines. Any lightning arrester

to protect the system must therefore be operative again immediately after the discharge, since very often a second and a third discharge follows immediately after the discharge within a second or less.

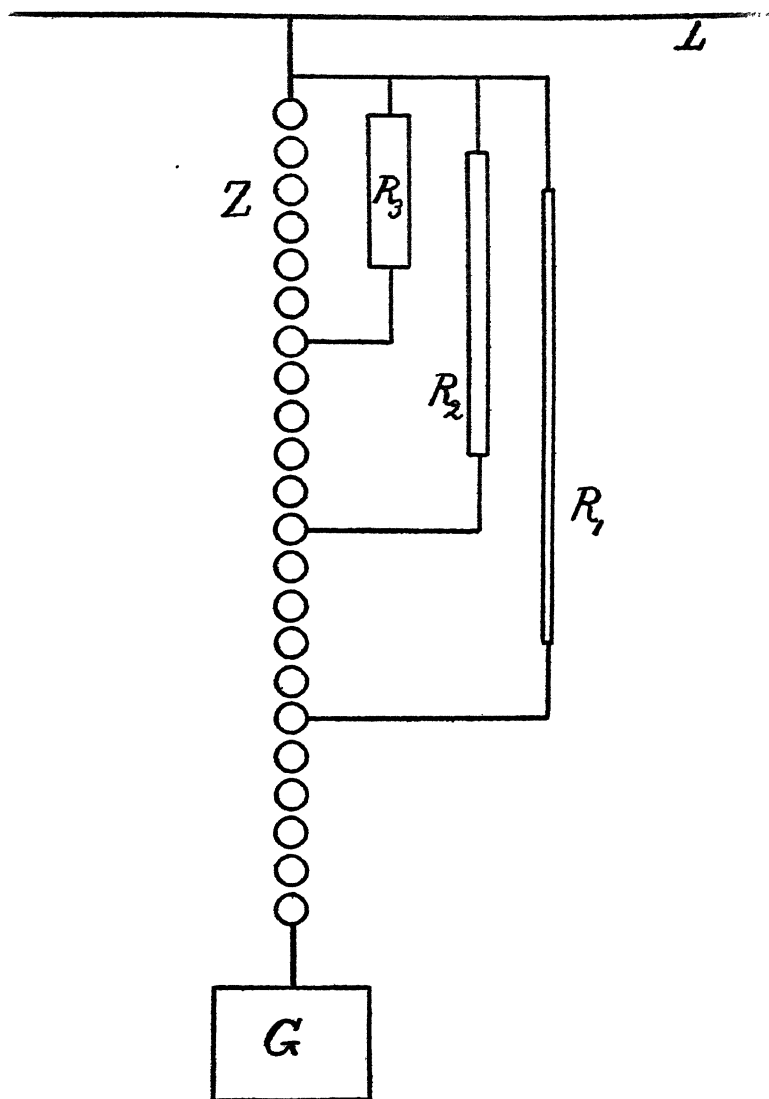


Fig. 27

Continuous discharges, or recurrent surges, (lightning lasting continuously for long periods of time with thousands of high voltage peaks per second), mainly originate in the circuits: by an arcing ground, spark discharge over broken insulators, faults in cables, etc. These phenomena, which have made their appearance only with the development of the modern high power high voltage electric systems, become of increasing severity and danger with the increase in size and power of electric systems.

Single strokes and multiple strokes, that is, all the disturbances due to atmospheric electricity, as cloud lightning, are safely taken care of by the modern multi-gap lightning arrester. In its usual form for high alternating voltages, it comprises a large number of spark gaps, connected between line and ground, and shunted by resistances of different sizes, as shown in Fig. 27, in such manner that a high pressure discharge of very low quantity, as the gradual accumulation of static charge on the system, discharges over a path of very high resistance R_1 , and so discharges inappreciably and even frequently invisibly. A disturbance of somewhat higher power finds a discharge path of moderate resistance R_2 , and so discharges with moderate current, that is, without shock on the system; while a high power disturbance finds a discharge path over a low resistance R_3 , and, if of very great power, even over a path of zero resistance, Z . On lower voltage, commonly only two resistances are used, one high and one moderately low, as shown by the diagram of a 2000 volt multi-gap arrester, Fig. 28.

The resistance of the discharge path of the present multi-gap arrester therefore is approximately inversely proportional to the volume of the discharge. This is an essential and important feature. Occasionally discharges of such large volume

occur, as to require a discharge path of no resistance, as any resistance would not allow a sufficient discharge to keep the voltage within safe limits. At the same time the discharge should not occur over a path without a resistance or of very low resistance, except when necessary, since the momentary short

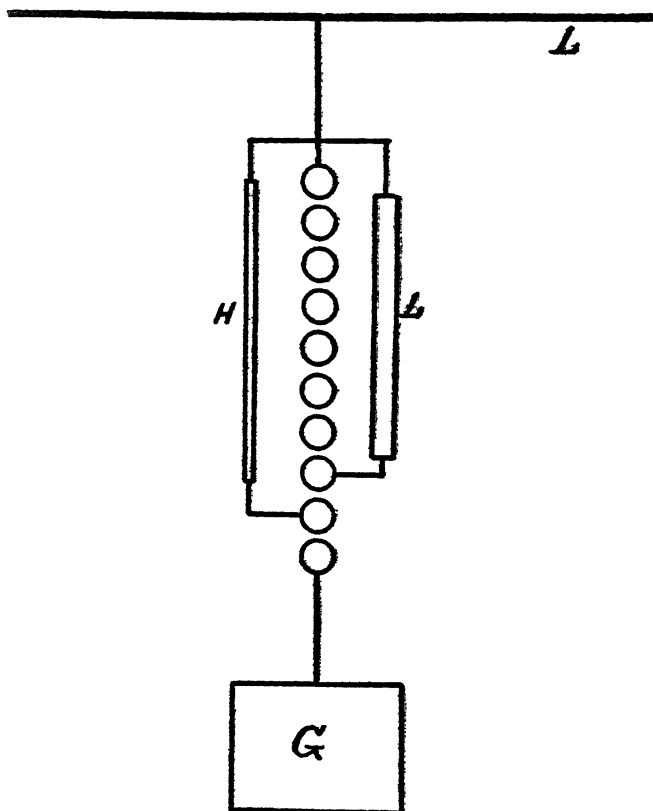


Fig. 28

circuit—that is, the short circuit for a part of the half wave—of a resistanceless discharge is a severe shock on the system, which must be avoided wherever permissible.

This type of lightning arrester takes care of single discharges and of multiple discharges, no matter how frequently

they occur or how rapidly they follow each other, with the minimum possible shock on the system. It cannot take care, however, of continuous lightning—those disturbances, mainly originating in the system, where the voltage remains excessive continuously (or rather rises thousands of times per second to excessive values), and for long times. With such a recurring surge, the multi-gap arrester would discharge continuously in protecting the system, until it destroys itself by the excessive power of the continuously succeeding discharges.

Where such continuous lightning may occur frequently, as in large high power systems, and the system requires protection against them, a type of lightning arrester which can discharge continuously, at least for a considerable time, without self-destruction, is necessary. The only lightning arrester which is capable of doing this, is the electrolytic, or aluminum arrester. In its usual form (cone or disc type) it comprises a series of cone-shaped aluminum cells, connected between line and ground through a spark gap. As soon as the voltage of the system rises above normal, by the value for which the spark gap is set, a discharge takes place through the aluminum cells, over a path of practically no resistance; but the volume of the discharge which passes, is not that given by the voltage on the system, but is merely that due to the excess voltage over the normal, since the normal voltage is held back by the counter e. m. f. of the aluminum cells. As a result—with strokes following each other, thousands per second, that is, with a recurrent surge—the aluminum arrester discharges continuously; but it can stand the continuous discharge for half an hour or more without damage, since it does not carry the short circuit current of the system, but merely the short circuit current of the excess voltage, and so protects the circuit

against continuous lightning for a sufficiently long time, until the cause of the high voltage can be found and eliminated.

Even the cone type of aluminum arrester discharges with a slight shock on the system, as the voltage must rise to the value of the spark gap, before the discharge begins, and in systems, in which even a small voltage shock is objectionable, as mainly in large underground cable systems, and also in cases where it is necessary to take care of recurrent surges for an indefinite time, the no-gap aluminum arrester becomes necessary. In principle, this type is the same as the cone type, but the aluminum cells are connected between the conductors and the ground without any spark gap, that is, are continuously in circuit, taking a small current. For this reason, the cells are made larger, and of different construction, so as to radiate the heat of the current which, while small, would still give a harmful temperature rise when allowed to accumulate. Being continuously in circuit, a no-gap aluminum arrester allows no sudden voltage rise whatever, however small it may be, that is, it acts just like a flywheel on the engine: while it allows gradual changes of voltages, any sudden change of voltage is anticipated and cut off, just as any sudden change of speed by the flywheel. The no-gap aluminum cell so can hardly be called a lightning arrester, but rather fulfills the duty of a shock absorber, an electrical flywheel on the voltage of the system, and as such finds its proper place on the bus bars of the station or substation, as "surge protector."

The three types of apparatus: the no-gap aluminum cell, the aluminum cone arrester, and the multi-gap lightning arrester, then are not different types of apparatus intended for the same purpose, but their operation and proper field of usefulness is different: the multi-gap arrester protects the system against atmospheric lightning and similar phenomena; the

aluminum cone arrester adds hereto protection against recurrent surges, where such surges may occur and the system requires protection against them, and thus finds its field, but at the same time requiring somewhat more attention than the multi-gap arrester; and the no-gap aluminum cell should be installed as electrical flywheel at the bus bars of the station, and in cable systems, usually in addition to other protection on lines and feeders; it requires, however, occasional attention, and continuously consumes a small amount of power.

Of other forms of lightning arresters, the magnetic blow-out 500 volt railway arrester is still in use to a large extent, but is beginning to be superseded by the aluminum cell. The multi-gap, being based on the non-arcing or rectifying property of the metal cylinders which exists only with alternating current, is not suitable for direct current circuits. In arc light circuits, that is, constant current circuits, horn gap arresters with series resistance are generally used, especially on direct current arc circuits, in which the multi-gap is not permissible. In such circuits of limited current, and very high inductance, the series resistance is not objectionable. Otherwise the horn gap arrester is still occasionally used outdoors as emergency arrester on transmission lines, set for a much higher discharge voltage than the station arrester, and then preferably without series resistance.

TWELFTH LECTURE

ELECTRIC RAILWAY

TRAIN CHARACTERISTICS

The performance of a railway consists of acceleration, motion and retardation, that is, starting, running and stopping.

The characteristics of the railway motor are:

1. Reliability.

2. Limited available space, which permits less margin in the design, so that the railway motor runs at a higher temperature, and has a shorter life, than other electrical apparatus. The rating of a railway motor is therefore entirely determined by its heating. That is, the rating of a railway motor is that output which it can carry without its temperature exceeding the danger limit. The highest possible efficiency is therefore aimed at, not so much for the purpose of saving a few percent. of power, but because the power lost produces heat and so reduces the motor output.

3. Very variable demands in speed. That is, the motor must give a wide range of torque and speed at high efficiency. This excludes from ordinary railway work the shunt motor and the induction motor.

The power consumed in acceleration usually is many times greater than when running at constant speed, and where acceleration is very frequent, as in rapid transit service, the efficiency of acceleration is therefore of foremost importance, while in cases of infrequent stops, as in long distance and inter-urban lines, the time of acceleration is so small a part of the total running time, that the power consumed during acceleration is a small part of the total power consumption, and high efficiency of acceleration is therefore of less importance.

Typical classes of railway service are :

1. Rapid transit, as elevated and subway roads in large cities.

Characteristics are high speeds and frequent stops.

2. City surface lines, that is, the ordinary trolley car in the streets of a city or town.

Moderate speeds, frequent stops, and running at variable speeds, and frequently even at very low speeds, are characteristic.

3. Suburban and interurban lines. That is, lines leading from cities into suburbs and to adjacent cities, through less densely populated districts.

Characteristics are less frequent stops, varying speeds, and the ability to run at fairly high speeds as well as low speeds.

4. Long distance and trunk line railroading.

Characteristics are: infrequent stops, high speeds, and a speed varying with the load, that is, with the profile of the road.

5. Special classes of service, as mountain roads and elevators.

Characteristics are fairly constant and usually moderate speed; a constant heavy load, so that the power of acceleration is not so much in excess of that of free running; and usually frequent stops. This is the class of work which can well be accomplished by a constant speed motor, as the three-phase induction motor.

The rate of acceleration and rate of retardation is limited only by the comfort of the passengers, which in this country permits as high values as 2 to $2\frac{1}{2}$ miles per hour per second, that is, during every second of acceleration, the speed increases at the rate of 2 to $2\frac{1}{2}$ miles per hour, so that one second after starting a speed of 2 to $2\frac{1}{2}$ miles per hour, 5 seconds

after starting a speed of 5×2 to $2\frac{1}{2} = 10$ to $12\frac{1}{2}$ miles per hour, etc., is reached.

Steam trains give accelerations of $\frac{1}{2}$ mile per hour per second and less with heavy trains, due to the lesser maximum power of the steam locomotive.

SPEED TIME CURVES

In rapid transit, and all service where stops are so frequent that the power consumed during acceleration is a large part of the total power, the speed time curves are of foremost importance, that is, curves of the car run, plotted with the time as abscissae, and the speed as ordinate.

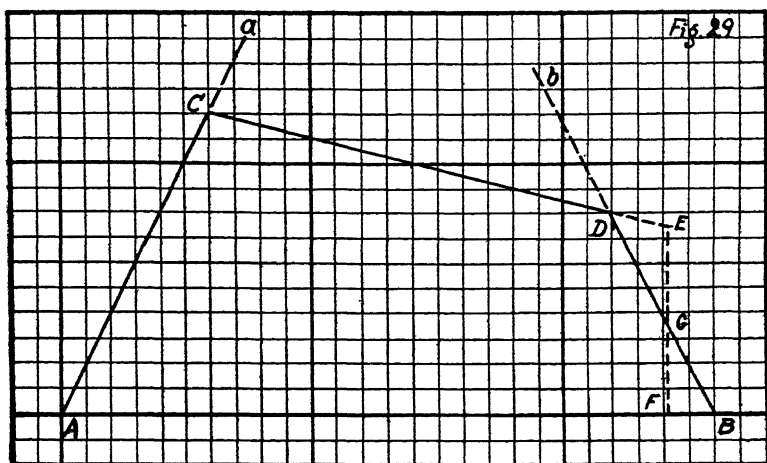


Fig. 29.

Choose for instance, a maximum acceleration and maximum braking of two miles per hour per second, and assuming a retardation of one-quarter mile per hour per second by friction (that is, assuming that the car slows down one-quarter mile per second, when running light on a level track); if then the time of one complete run between two stations is given equal to A B in Fig. 29, the simplest type of run consists of constant acceleration, from A to C, on the line A a, drawn

under a slope of two miles per hour per second; at C the power is shut off and the car coasts on the slope C D, of one-quarter mile per hour per second, until at D, where the coasting line cuts the braking line bB, (which also is drawn at the slope of two miles per hour per second), the brakes are applied and the car comes to rest, at B. As the distance traveled is speed times time, the area A C D B so represents the distance traveled, that is, the distance between the two stations, and all speed time curves of the same type therefore must give the same area. During acceleration, energy is put into the car, and stored by its momentum, which is proportional to the weight of the car and the square of the speed. It is therefore at a maximum at C. A part of the energy represented by the car speed is consumed during coasting in overcoming the friction; the rest is destroyed by the brakes. Assuming, as approximation, constant friction, the energy consumed by the car friction on the track, for runs of the same distance, is constant, and the energy destroyed by the brakes is represented by the speed at the point B, where the brakes are applied. The lower therefore this point B is, the less power is destroyed by the brakes, and the more efficient is the run. More accurately, by prolonging C D to E so that area D E G = B F G, the area A C E F also is the distance between the stations, and E F so would be the speed at which the car arrives at the next station, if no brakes were applied, and the energy corresponding thereto has to be destroyed by the brakes; that is, represents the energy lost during the run, and should be made as small as possible, to secure efficiency.

The ratio of the energy used for carrying the car across the distance between the stations—that is, energy consumed by track friction, (plus energy consumed in climbing grades, where such exist) to the total energy input, that is, track fric-

tion plus energy consumed in the brakes, is the *operation efficiency* of the run.

As an illustration, a number of such runs, for constant time of the run, of 130 seconds, and constant distance between the stations, that is, constant area of the speed time diagram, are plotted in Figs. 29 to 37.

1. Constant acceleration of two miles per hour per second, coasting at one-quarter mile per hour per second, and braking at two miles per hour per second. Here the energy consumed by the brakes is given by the speed $E F = 34.5$ miles per hour, while the maximum speed reached is 60 miles per hour.

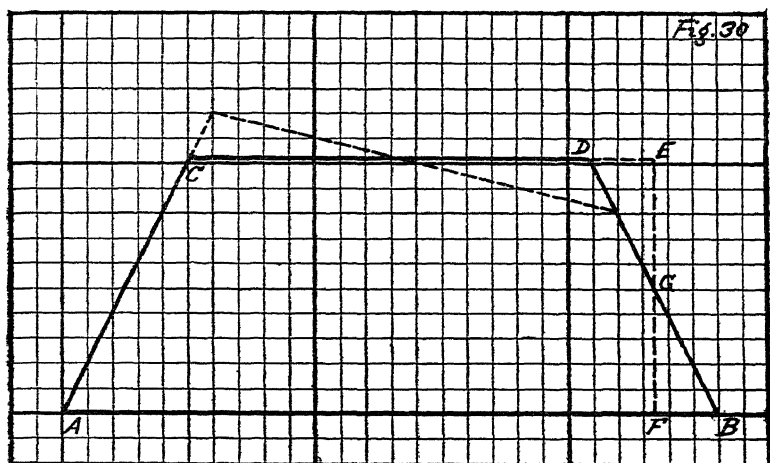


Fig. 30.

2. Acceleration and retardation at two miles per hour per second. Constant speed running between. Fig. 30. Compared with 1, (which is shown in 30 in dotted lines), the maximum speed is slightly reduced, e. g., to 51 miles per hour, but the speed of application of the brakes, and therefore the energy lost in the brakes, is increased. That is, running at constant speed, between acceleration and braking, is less efficient than coasting

with decreasing speed. Besides this, at the low power required for constant speed running, the motor efficiency usually is already lower. It therefore is uneconomical to keep the power on the motors after acceleration, and more economical to continue to accelerate until a sufficient speed is reached to coast until the brakes have to be applied for the next station. Obviously, this is not possible where the distance between the stations is so great, that in coasting the speed would decrease too much to make the time, and so applies only to the case of runs with frequent stops, as rapid transit.

3. Constant acceleration of one mile per hour per second, braking at two miles, coasting one-quarter mile. Dia-

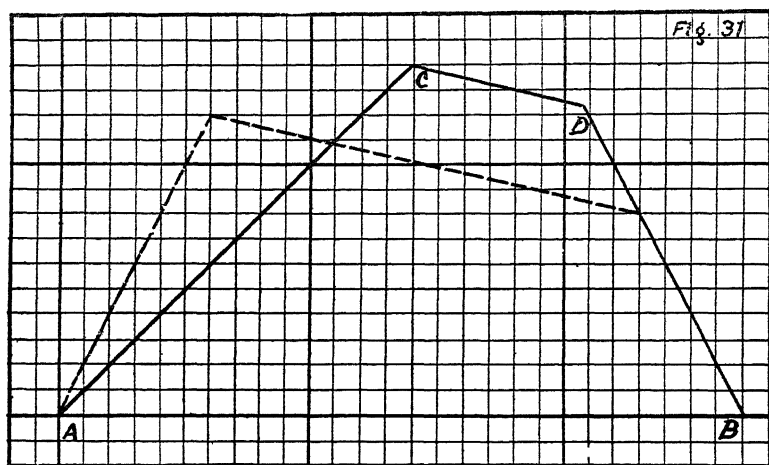


Fig. 31.

gram 1 is shown in the same figure 31, for comparison. As seen, with the lower rate of acceleration, the maximum speed is greater, and the lost speed, or speed E F, which is destroyed by the brakes, is greater, that is, the efficiency of the run is lower.

4. Constant acceleration and braking of one mile per hour per second, coasting at one-quarter mile. In this case,

the run between the stations cannot be made in 130 seconds. For comparison, τ is shown dotted in Fig. 32. Here the maxi-

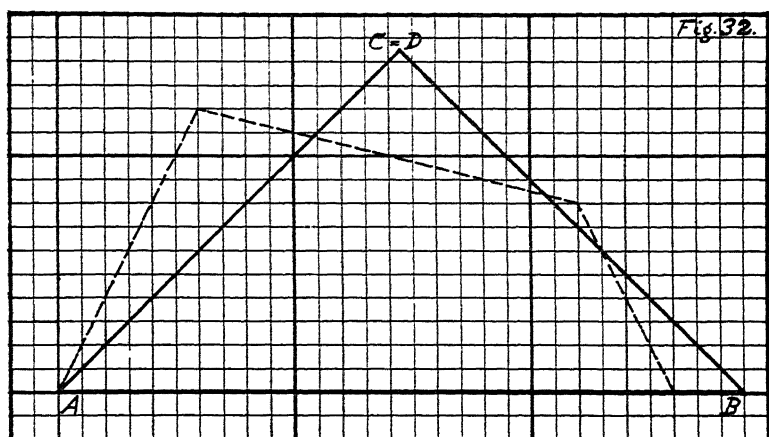


Fig. 32.

mum speed and the lost speed are still greater, that is, the efficiency of the run still lower, and at least 145 seconds are required. That is, the higher the rate of acceleration and of braking, the less is the maximum speed required, and the higher the operation efficiency. With constant acceleration up to the maximum speed, the operation therefore is the more efficient the higher a rate of acceleration and of braking is used. While very rapid acceleration requires more power developed by the motor and put into the car, the time during which the power is developed is so much shorter, that the energy put into the car, or power times time of power application, is less than with the lower rate of acceleration.

The highest operation efficiency, in the case of frequent stops, therefore is produced by constant acceleration at the highest permissible rate, coasting without power, and then braking at the highest permissible rate, as given by τ .

During acceleration at constant rate, from A to C, the motor however runs on the rheostat. That is, at all speeds below the maximum, to produce the same pull as at the maximum speed C, the motor consumes the same current and so the same power; while the power which it puts into the train is proportional to the speed, and therefore is very low at low speeds. Or in other words, the motor during constant acceleration, consumes power corresponding to maximum speed, while the useful power corresponds to the average speed, which during A C is only half the maximum; and so only half the available power is put into the car, the other half being wasted in the resistance, and the *motor efficiency* during constant acceleration therefore must be less than 50%.

Constant acceleration up to maximum speed, while giving the best operation efficiency, so gives a very poor motor efficiency and thereby low total efficiency, (the total efficiency being the ratio of the useful energy to the total energy put into the motors, that is, is operation efficiency times motor efficiency).

This is the arrangement necessary for a constant speed motor, as the induction motor; but it does not give the best total efficiency, but a better total efficiency is produced by accelerating partly on the motor curve, that is, at a decreasing rate. This sacrifices some operation efficiency, but increases the motor efficiency greatly, and so, if not carried too far, increases the total efficiency.

The speed time curves of the motor are shown in Fig. 33, and the current consumption is also plotted in this figure. Acceleration is constant from A to M, on the rheostat, and at constant current consumption, from M, onwards, the acceleration decreases, first slightly, then faster, but the current also decreases, first rapidly, and then more slowly; and the

efficiency, plotted in Fig. 33, rises from 0% at A, to 90% at M, and then remains approximately constant, while the speed still increases.

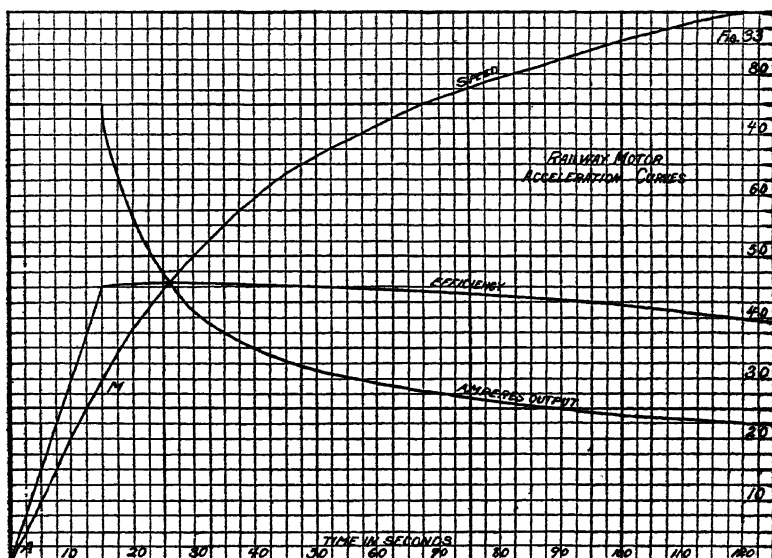


Fig. 33.

6. This gives the speed time curve of the car, Fig. 34, with acceleration on the motor curve and with maximum values of acceleration and braking 2, the coasting value one-quarter; that is, the same as 1, and 1 is shown in dotted lines in the same figure. The acceleration is constant, on the rheostat, from A to M; at M the rheostat is cut out, and the acceleration continues on the motor curve, at a gradually decreasing rate, until at C the power is shut off and the car coasts until the brakes are applied. The area A M C D B, representing the distance between the stations, is the same as in 1; the operation efficiency is somewhat lower, but the total current consumption, as shown by the curves of current, shown together

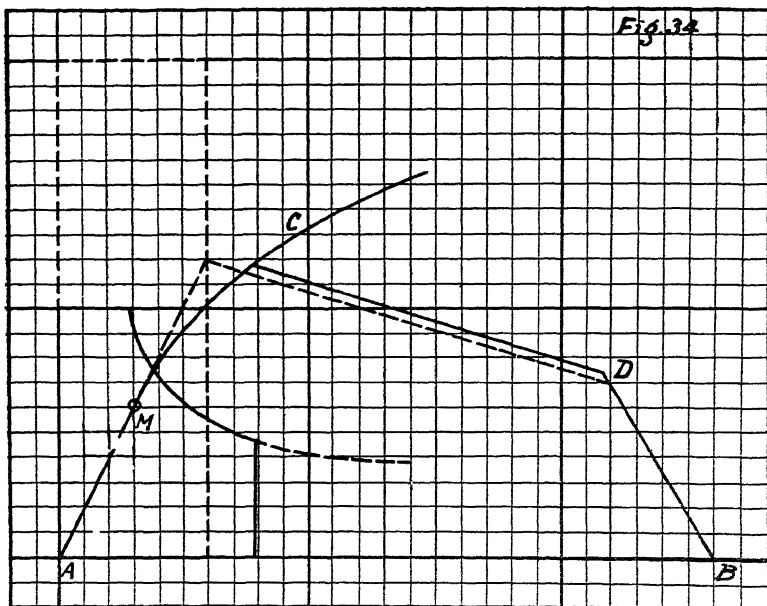


Fig. 34.

with the speed time curves, is much less, and the power consumption therefore is less; that is, the total efficiency is higher.

7. Fig. 35 gives another speed time curve in which, however, the motor is geared for too low a speed; so the motor

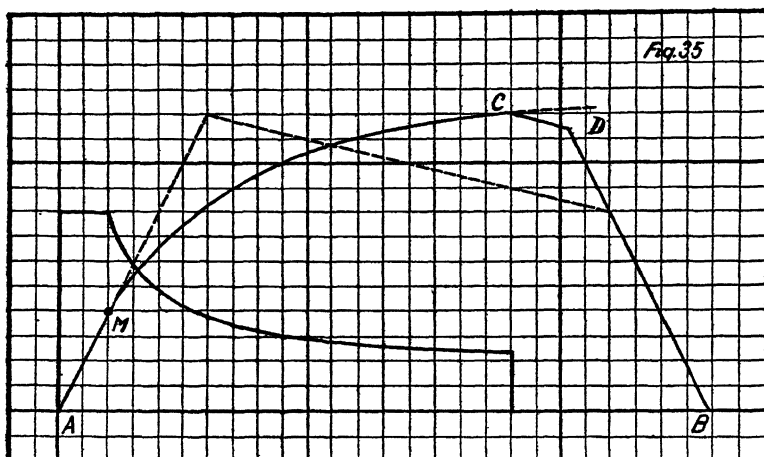


Fig. 35.

curve is reached too early, and the power has to be kept on for too long a time, to make the run in time. As seen from the current curves, here the loss in car efficiency by the decreased acceleration on the motor curve is greater than the saving in motor efficiency, and the power consumption by the motor is greater than that without running on the motor curve.

That is, the total efficiency of operation is increased by doing some of the accelerating on the motor curve, but may

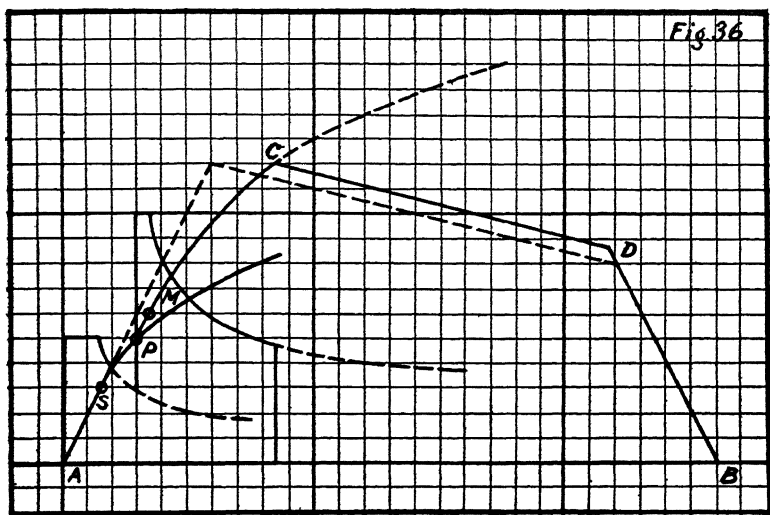


Fig. 36.

be impaired again by carrying this too far. Usually the rheostat is all cut out and the acceleration continues on the motor curve, from about half speed onwards.

8. During the first half of the acceleration on the rheostat, when more than half the voltage is consumed in the rheostat, half the current can be saved by connecting two motors in series; that is, by series parallel control on the motors, as shown in Fig. 36. If, however, the series connection of motors is maintained too long, as shown in Fig. 37,

increased power consumption in the brakes, by the higher motor efficiency.

CONCLUSION

In short distance runs the efficiency is highest in running on series parallel control as much as possible on the motor curve, with as high a rate of average acceleration and retardation as possible, and coasting between acceleration and retardation; that is, not keeping the power on longer than necessary.

The longer the distance, the less important is high rate of acceleration and retardation, and for long distance running the rate of acceleration and retardation is of little importance.

Therefore speed time curves are specially important in rapid transit service, and in general, in running with frequent stops.

The heating of the motor at high acceleration, that is, with large current, is less than with low acceleration, that is, smaller current, because the current is on a much shorter time.

Feeding back in the line by using the motors as generators is rarely used; because with an efficient speed time curve, using coasting, the speed when putting on the brakes is already so low that usually not enough power can be saved to compensate for the complication and the increased heating of the motors, when carrying current also in stopping. The motors are occasionally used as brakes, operating as generators on the rheostat. This, however, puts an additional heating on the motors; and is therefore not much used in this country, where the highest speed which the motor equipment can give is desired.

With induction motors, feeding back in the line is simplest, because induction motors become generators above

synchronism, and so feed back when running down a long hill. Therefore on mountain railways, induction motors have the advantage.

In an induction motor there is no running on the motor curve, and so the efficiency of acceleration is lower.

Objection to the series motor is the unlimited speed; that is, when running light, it runs away. In railroading this is no objection, because the motor is never running light and somebody is always in control.

In elevator work the series motor is objectionable, due to the unlimited speed; therefore a limited speed motor is necessary. In elevators frequent stops, and so efficient acceleration are necessary; therefore a compound motor is best, that is, a motor having a shunt field to limit the speed and a series field (which is cut out after starting) to give efficient acceleration.

THIRTEENTH LECTURE

ELECTRIC RAILWAY: MOTOR CHARACTERISTICS

THE economy of operation of a railway system, station, lines, etc., decreases, and the amount of apparatus, line copper, etc., which is required, increases with increasing fluctuations of load; the best economy of an electric system therefore requires as small a power fluctuation as possible.

The pull required of the railway motor during acceleration, on heavy grades, etc., is, however, many times greater than in free running. In a constant speed motor, as a direct current shunt motor or an alternating current induction motor, the power consumption is approximately proportional to the torque of the motor and thus to the draw bar pull that is given by it. With such motors, the fluctuation of power consumption would thus be as great as the fluctuation of pull required. In a varying speed motor, as the series motor, the pull increases with decreasing speed; and the power consumption, which is approximately proportional to pull times speed, varies less than the pull of the motor. The fluctuation of load produced in the circuit by a series motor therefore is far less than that produced by a shunt or induction motor—the former economizing power at high pull by a decrease of speed; the series motor thus gives a more economical utilization of apparatus and lines than the shunt or induction motor, and is therefore almost exclusively used.

The torque, and so the pull produced by a motor, is approximately proportional to the field magnetism and the armature current; that is, neglecting the losses in the motor, or assuming 100% efficiency, the torque is proportional to the product of magnetic field strength and armature current.

In a shunt motor, at constant supply voltage e , the field exciting current, and thus the field strength, is constant; and the torque, when neglecting losses, is thus proportional to the armature current, as shown by the curve T_0 in Fig. 38. From

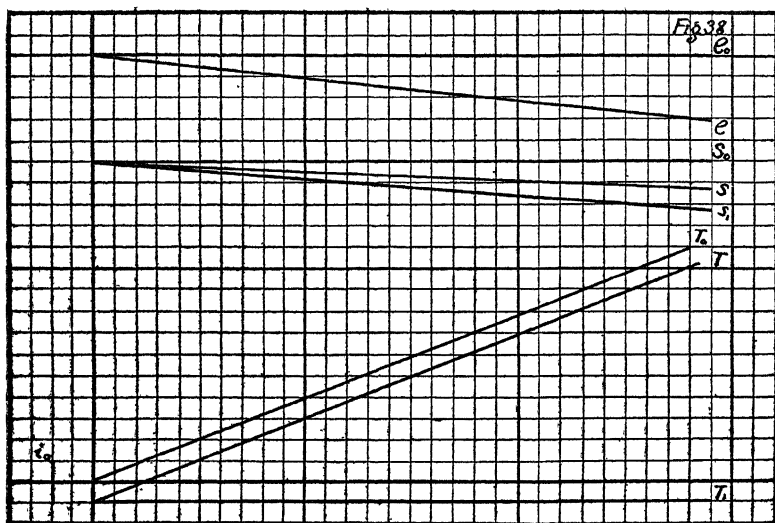


Fig. 38.

this torque is subtracted the torque consumed by friction losses, coreloss, etc. (which, at approximately constant speed and field strength, is approximately constant and is shown by the curve T_1) thus giving as net torque of the motor, the curve T . Neglecting losses, the speed of the motor would be constant, as given by line S_0 ; since at constant field strength, to consume the same supply voltage e_0 , the armature has to revolve at the same speed. As, however, with increasing load and therefore increasing current, the voltage available for the rotation of the armature decreases by the ir drop in the armature, as shown by the curve e at constant field strength, the speed decreases in the same proportion, as shown by the curve S_1 . The field strength, however, does not remain perfectly constant, but with

increasing load the field magnetism slightly changes: it decreases by field distortion and demagnetization, and the speed therefore increases in the same proportion, to the curve *S*. The current used as abscissae in Fig. 38 is the armature current. The total current consumed by the motor is, however, slightly greater, namely, by the exciting current i_0 ; and, plotted for the total current of the motor as abscissae, all the curves in Fig. 38 are therefore shifted to the right, by the amount of i_0 , as shown in Fig. 39.

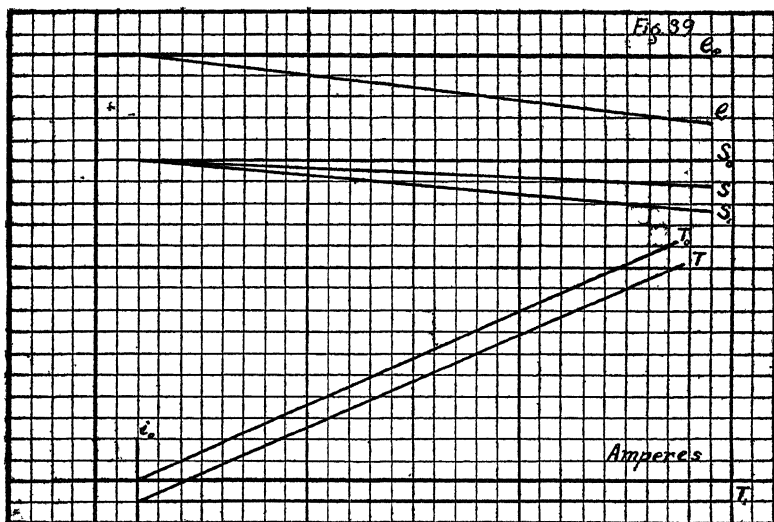


Fig 39.

If in the shunt motor, the supply voltage changes, the field strength, which depends upon the supply voltage, also changes; it decreases with a decrease of the supply voltage, and the current required to produce the same torque therefore increases in the same proportion. If the magnetic field is below saturation, the field strength decreases in proportion to the decrease of supply voltage, and the current thus increases in proportion to the decrease of supply voltage, while the speed re-

mains the same, the armature produces the lower voltage by revolving in the lower field at the same speed. If the magnetic field is highly over-saturated and does not therefore appreciably change with a moderate change of supply voltage and so of field current, the armature current required to produce the same torque also does not appreciably change with a moderate drop of supply voltage, but the speed decreases, since the armature must now consume less voltage in the same field strength.

Depending on the magnetic saturation of the field: with a decrease of the supply voltage the current consumed by the shunt motor to produce the same torque, therefore increases the more, the lower the saturation, and the speed decreases the more, the higher the saturation.

In general, a drop of voltage in the resistance of lines and feeders does not much affect the speed of the shunt motor, but increases the current consumption, thus still further increasing the drop of voltage; so that in a shunt motor system, lines and feeders must be designed for a lower drop in voltage than is permissible for a series motor.

The three-phase induction motor in its characteristics corresponds to a shunt motor with under-saturated field, except that the effect of a drop of voltage is still more severe; as not only the amount, but usually the lag of current also increases, thus causing more drop in voltage; and the maximum torque of the motor is limited, and decreases with the square of the voltage. Hence, while in a series motor system the lines and feeders are designed for the average load or average voltage drop (and practically no limit exists to the permissible maximum voltage drop), with an induction motor, the maximum permissible voltage drop is limited by the danger of stalling the motors.

In the series motor, the armature current passes through the field, and with increasing load and thus increasing current, the field strength also increases; the torque of the motor therefore increases in a greater proportion than the current. Neg-

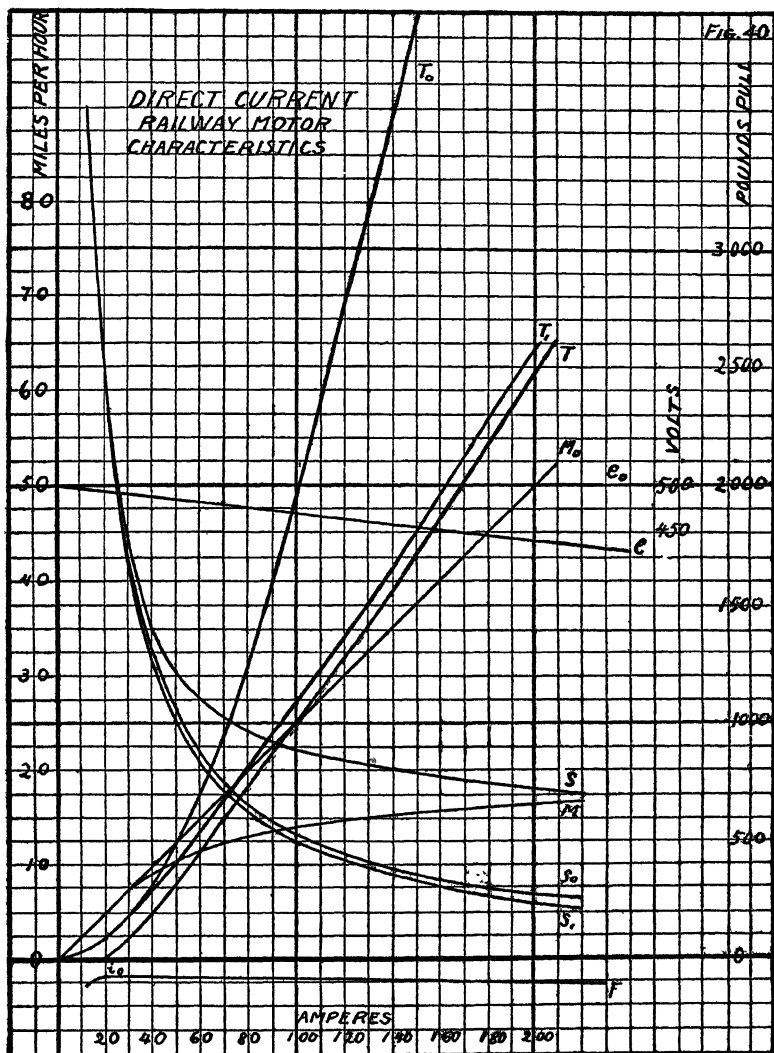


Fig. 40.

lecting losses and saturation, the field strength is proportional to the current; the torque being proportional to the current times field strength, therefore is proportional to the square of the current, as shown by the curve T_0 in Fig. 40. The supply voltage, however, has no direct effect on the torque; but with the same current consumption, the motor gives the same torque, regardless of the supply voltage. The speed, at constant supply voltage, changes with the field strength and thus with the current: the higher the field strength, the lower is the speed at which the armature consumes the voltage. Since the field strength—neglecting losses and saturation—is proportional to the current, the speed of the series motor would be inversely proportional to the current, as shown by the curve S_0 in Fig. 40.

As the voltage available for the armature rotation decreases with increasing current, from e_0 to e , by the ir drop in the field and armature, the speed decreases in the same proportion, from the curve S_0 to the curve S_1 .

In reality, however, the field strength, as shown by the curve M_0 , is proportional to the current only at low currents; but for higher currents the field strength drops below, by magnetic saturation, as shown by the curve M ; and ultimately, at very high currents, it becomes nearly constant. In the same ratio as the field strength drops below proportionality with the current, the speed increases and the torque decreases. The actual speed curve is therefore derived from the curve S_1 by increasing the values of the curve S_1 in the proportion, M_0 to M , and is given by the curve S ; and in the same proportion the torque is decreased to the curve T_1 . From this torque curve the lost torque is now subtracted; that is, the torque representing the power consumed in friction and gear losses, hysteresis and eddy currents, etc. Some of the losses of power are

approximately constant; others are approximately proportional to the square of the current; and the lost torque, being equal to the power loss divided by the speed, can therefore be assumed as approximately constant: somewhat higher at low and high speeds, as shown by curve F. The net torque then is given by the curve T. As seen, it is approximately a straight line, passing through a point i_0 , which is the "running light current," and its corresponding speed, the "free running speed" of the motor. At this current i_0 , the speed is highest; with increase of current it drops first very rapidly, and then more slowly; and the higher the saturation of the motor field is, the slower becomes the drop of speed at high currents.

The single-phase alternating current motors are either directly or inductively series motors, and so give the same general characteristics as the direct current series motor. In the alternating current motors, however, in addition to the ir drop an ix drop exists; that is, in addition to the voltage consumed by the resistance, still further voltage is consumed by self-induction; and the voltage e available for the armature rotation thus drops still further, as seen in Fig. 41. Since the self-induction consumes voltage in quadrature with the current, the inductive drop is not proportional to the current, but is small at low currents, and greater at high currents; e therefore is not a straight line, but curves downwards at higher currents. The speed, S_1 , is dropped still further by the inductive drop of voltage, to the curve S_1 , and then raised to the curve S by saturation. The effect of saturation in the alternating current motor usually is far less, since the magnetic field is alternating, and good power factor requires a low field excitation, and therefore high saturation cannot well be reached. The torque curves are the same as in the direct current motor, except that the effect of saturation is less marked.

In efficiency, the shunt or induction motor, and the series motor are about equal; and both give high values of efficiency over a wide range of current. A wide range of current, however, represents a wide range of speed in the series motor, and

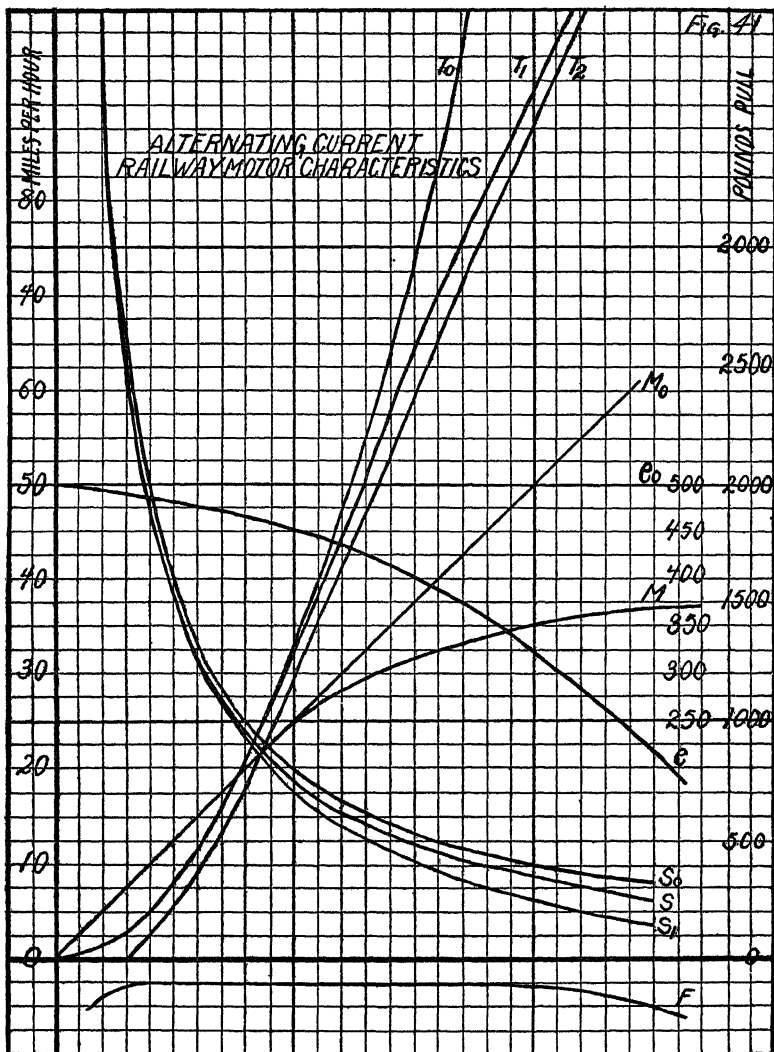


Fig. 41.

nearly constant speed in the shunt motor; therefore while the series motor can operate at high efficiency over a wide range of speed, the shunt motor shows high efficiency only at its proper speed.

In regard to the effect of a change of supply voltage, as is caused, for instance, by a drop of voltage in feeders and mains, the series motor reacts on a change of voltage by a corresponding change of speed, but without change of current; while the shunt motor and induction motor react on a change of supply voltage by a change of current, with little or no change of speed. As the limitation of a system usually is the current, at excessive overloads on the system, resulting in heavy voltage drop, the series motors run slower, but continue to move; while the induction motor is liable to be stalled.

FOURTEENTH LECTURE

ALTERNATING CURRENT RAILWAY MOTOR.

IN a direct current motor, whether a shunt or a series motor, the motor still revolves in the same direction, if the impressed e. m. f. be reversed, as field and armature both reverse. Since a reversal of voltage does not change the operation of the motor, such a direct current motor therefore can operate also on alternating current. With an alternating voltage supply, the field magnetism of the motor also alternates; the motor field must therefore be laminated, to avoid excessive energy losses and heating by eddy currents (currents produced in the field iron by the alternation of the magnetism) just as in the direct current motor the armature must be laminated.

In the shunt motor—in which the supply current divides between field and armature—when built for alternating voltage, arrangements must be made to have the current in the field (or rather the field magnetism) and the current in the armature, reverse simultaneously. In the series motor, in which the same current traverses field and armature, the field magnetism and the armature current are necessarily in phase with each other, or nearly so. Only the series or varying speed type of alternating current commutator motor has so far become of industrial importance.

In the alternating current motor in addition to the voltage consumed by the resistance of the motor circuit and that consumed by the armature rotation, voltage is also consumed by self-induction; that is, by the alternation of the magnetism. The voltage consumed by the resistance represents loss of power, and heating, and is made as small as possible in any

motor. The voltage consumed by the rotation of the armature, or "e. m. f. of rotation," is that doing the useful work of the motor, and so is an energy voltage, or voltage in phase with the current; just as the voltage consumed by the resistance is in phase with the current. The voltage consumed by self-induction, due to the alternation of the magnetism, or "e. m. f. of alternation", is in quadrature with the current, or wattless; that is, it consumes no power, but causes the current to lag, and so lowers the power factor of the motor; that is, causes the motor to take more volt-amperes than corresponds to its output, and so is objectionable.

The useful voltage, or e. m. f. of rotation of the motor, is proportional to the speed; or rather the "frequency of rotation", N_0 , is proportional to the field strength F , and to the number of armature turns m . The wattless voltage, or self-induction of the field, is proportional to the frequency N , to the field strength F , and the number of field turns n . The ratio of the useful voltage to the wattless voltage therefore is $mN_0 \div nN$, and to make the useful voltage high and the wattless voltage low, therefore requires as high a frequency of rotation N_0 and as low a frequency of supply N , as possible. Thus the commutator motors of more than 25 cycles give poor power factors; and for a given number of revolutions N_0 , which is number of revolutions per second times number of pairs of poles, therefore is the higher, the more poles the motor has. Hence a greater number of poles are generally used in an alternating current than in a direct current motor.

Good direct current motor design requires a strong field and weak armature, to get little field distortion and therefore good commutation; that is high n and low m . But such proportions, even at low supply frequency N and high frequency of rotation N_0 , would give a hopelessly bad power factor, and

thus a commercially impractical motor. In the alternating current commutator motor, it is therefore essential to use as strong an armature and as weak a field (that is, as large a number of armature turns m and as low a number of field turns n) as possible. Very soon, however, a limit is reached in this direction, even if the greater field distortion and the resultant bad commutation were not to be considered: the armature also has a self-induction; that is, the alternating magnetism produced by the current in the armature turns consumes a wattless e. m. f. This magnetism is small in a direct current motor, but with many armature turns and few field turns it becomes quite considerable; and so, while a further decrease of the field turns and increase of the armature turns reduces the self-induction of the field—which varies with the square of the field turns—it increases the self-induction of the armature—which varies with the square of the armature turns. There is thus a best proportion between armature turns and field turns, which gives the lowest total self-induction. This is about in this proportion: armature turns m to field turns $n = 2 \div 1$; and at this proportion the power factor of the motor, especially at low and moderate speeds, is still very poor.

In alternating current commutator motors it is therefore essential to apply means to neutralize the armature self-induction and armature reaction, so as to be able to increase the proportion of armature turns to field turns sufficiently to get good power factors. This is done by surrounding the armature with a stationary “compensating winding” closely adjacent to the armature conductors, located in the field pole faces, and traversed by a current opposite in direction to the current in the armature, and of the same number of ampere turns; so that the armature ampere turns and the ampere turns of the compensating winding neutralize each other, and the armature

reaction, that is, the magnetic flux produced by the armature current, and the self-induction caused by it, disappear.

This compensating winding for neutralizing the armature self-induction was introduced by R. Eickemeyer in the early days of the alternating current commutator motor, and since then all alternating current commutator motors have it; so that the electric circuits of all alternating current commutator motors comprise an armature winding A, a field winding F, and a compensating winding C.

Since the compensating winding cannot be identically at the same place as the armature winding (the one being located in slots in the pole faces, the other in slots in the armature face) there still exists a small magnetic flux produced by the armature winding: the "leakage flux", analogous to the leakage flux of the induction motor; and the number of armature turns cannot be increased indefinitely, otherwise the armature self-induction, due to this leakage flux, would become appreciable, and the power factor would decrease again. The minimum total self-induction of the motor with compensating winding occurs at a number of armature turns equal to 3 to 5 times the field turns; at this proportion, the power factor is already very good at low speeds, and the motor is industrially satisfactory in this regard.

For best results, that is, complete compensation and therefore zero magnetic field of armature reaction, it is, however, necessary not only to have the same number of ampere turns in the compensating winding as on the armature, but also to have these ampere turns distributed in the same manner around the circumference. With the usual armature winding this is not the case, but the armature conductors cover the whole circumference; while the compensating coil conductors cover only the pole arc, as the space between the poles is taken up by the

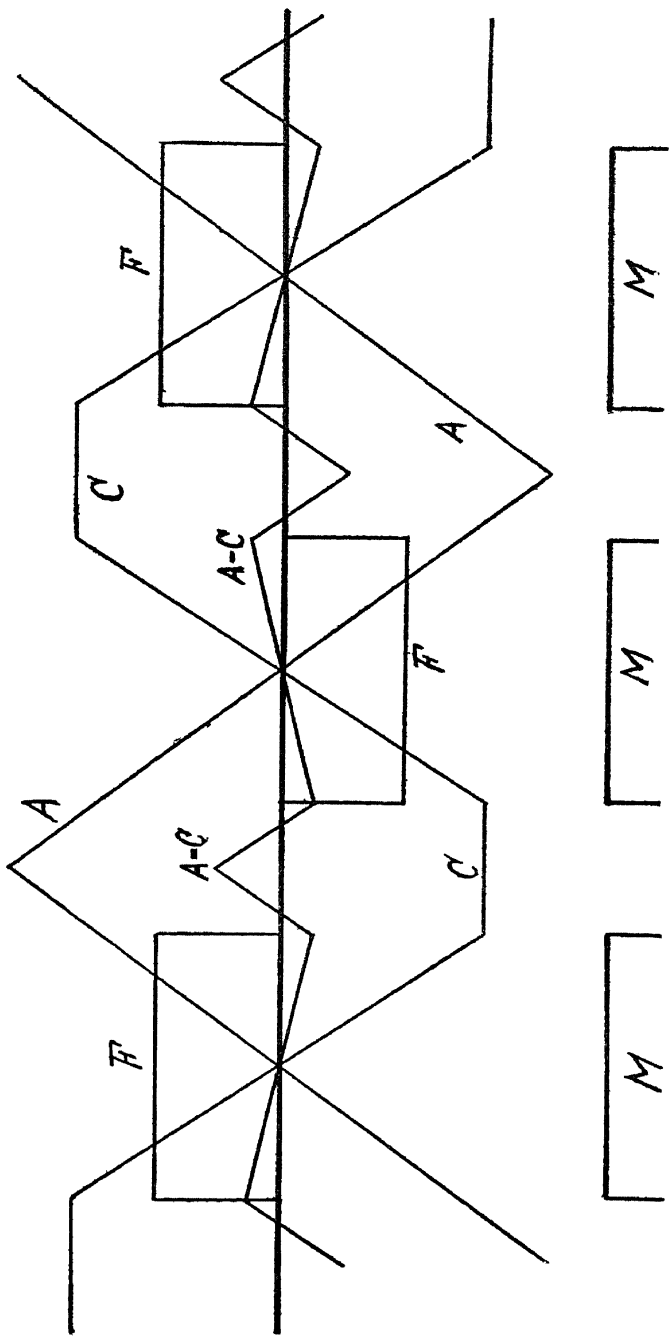


Fig. 42

field winding. That is, the magnetic distribution around the armature circumference is as shown developed in Fig. 42: the field gives a flat topped distribution, the armature a peaked, and the compensating winding has a small flat top and with the total ampere turns of the compensating winding equal to those of the armature, the compensating winding preponderates in front of the field poles, the armature between the field poles, or at the brushes, and there is thus a small magnetic field of armature reaction remaining at the brushes, just where it is objectionable for commutation.

As it is not feasible to distribute the compensating winding over the whole circumference of the stator, the armature winding is arranged so that its ampere turns cover only the pole arcs. This is done by using fractional pitch in the armature; that is, the spread of the armature coil or the space between its two conductors, is made, not equal to the pitch of the pole, as shown in Fig. 43, but only to the pitch of the pole arcs, as shown in Fig. 44. With such fractional pitch winding, the currents in the upper and the lower layer of the armature conductors, in the space between the poles, flow in opposite directions, and so neutralize, leaving only that part of the armature winding in front of the pole arcs as magnetizing. Hereby the distribution of the armature ampere turns is made the same as that of the compensating winding, and so complete compensation is realized.

The compensating winding may be energized by the main current, and so connected in series with the field and armature; or the compensating winding may be short circuited upon itself, and so energized by an induced current acting as a secondary of a transformer to the armature as primary; and as in a transformer, primary and secondary current have the same number of ampere turns (practically) and flow in opposite

directions, such "inductive compensation" is just as complete compensation as the "conductive compensation" produced by passing the main current through the compensating winding.

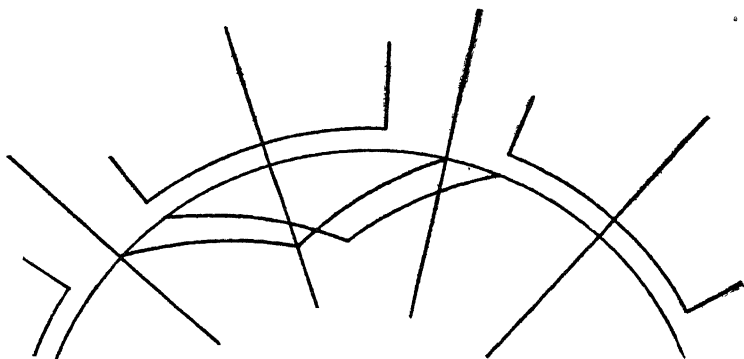


Fig. 43

Vice versa, the armature may be short circuited and so used as secondary of a transformer, with the compensating winding acting as primary. In either of these motor types,

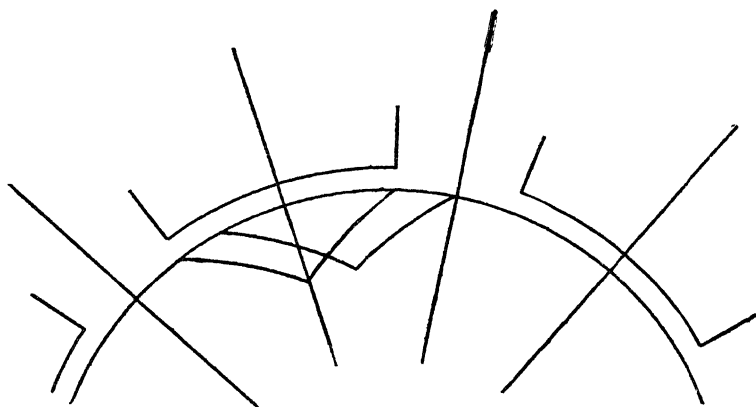


Fig. 44

which comprise primary and secondary circuits, that is, in which armature and compensating winding are not connected directly in series, but inductively, the field may be energized

by the primary or supply current, or by the secondary or induced current. In such a motor embodying a transformer feature, instead of impressing the supply voltage upon one circuit as primary, while the other is closed upon itself as secondary, the supply voltage may be divided in any proportion between primary and secondary.

As primary and secondary current of a transformer are proportional to each other, it is immaterial, regarding the variation of the current in the different circuits with the load and speed, whether the circuits are directly in series, or by transformation; that is, all these motors have the same speed—torque—current characteristics, as discussed in the preceding lecture, and differ only in secondary effects, mainly regarding commutation.

The use of the transformer feature also permits, without change of supply voltage, to get the effect of a changed supply voltage, or a changed number of field turns, by shifting a circuit over from primary to secondary or vice versa. For instance, if the armature is wound with half as many turns, that is, for half the voltage and twice the current, as the compensating winding, by changing the field from series connection with the compensating winding to series connection with the armature, the current in the field and thus the field strength, is doubled; that is, the same effect is produced as would be by doubling the number of field turns.

According to the relative connection of the three circuits, armature A, compensating circuit C, and field F, alternating current commutator motors of the series type can be divided into the classes shown diagrammatically in Fig. 45:

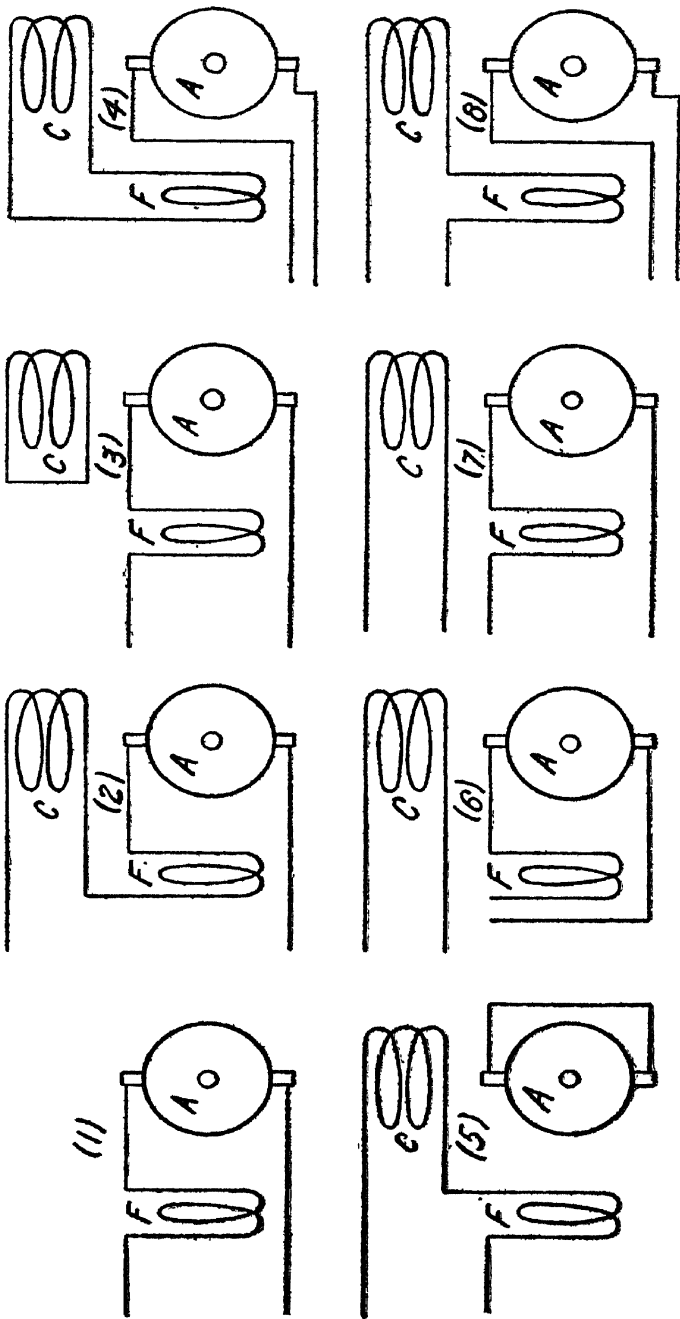


Fig. 45

Primary:	Secondary:	
$A + C + F$	———	Conductively Compensated Series Motor. (2).
$A + F$	C	Inductively Compensating Series Motor. (3).
A	$C + F$	Inductively Compensating Series Motor with Secondary Excitation, or Inverted Repulsion Motor. (4).
$C + F$	A	Repulsion Motor. (5).
C	$A + F$	Repulsion Motor with Secondary Excitation (6).
$C \& A + F$		Series Repulsion Motor A. (7).
$C + F \& A$		Series Repulsion Motor B. (8).

The main difference between these types of motors is found in their commutation.

In a direct current motor, with the brushes set at the neutral; that is, midway between the field poles (as is customary in a reversible motor like the series motor), the armature turn, which is short circuited under the brush during the commutation, encloses all the lines of magnetic force of the field; so during this moment it does not cut any lines of force by its rotation, and thus no e. m. f. is induced in this turn; that is, no current is produced, if the armature reaction is compensated for, or is otherwise negligible. If the motor has a considerable armature reaction, and thus a magnetic field at the brushes, this magnetic field of armature reaction induces an e. m. f. in the short circuited turn under the brush, and so

causes sparking. Hence high armature reaction impairs the commutation of the motor.

In an alternating current series motor the armature reaction is neutralized by the compensating winding, and therefore no magnetic field of armature reaction exists; hence no e. m. f. is induced in the turn short circuited under the brush by its rotation through the magnetic field. As this field, however, is alternating, an e. m. f. is induced in the short circuited turn by the alternations of the lines of magnetic force enclosed by it, and causes a short circuit current and in that way, sparking. This e. m. f., being due to the alternation of the enclosed field flux, is independent of the speed of rotation; it also exists with the motor at a standstill, and is a maximum in the armature turn under the brush, as this encloses the total field flux. The position of the armature turn during commutation, which in a direct current motor is the position of zero induced e. m. f., is therefore in an alternating current motor, the position of maximum induced e. m. f., but induced not by the rotation of the turn, but by the alternation of the magnetism. That is, this turn is in the position of a short circuited secondary to the field coil of the motor as primary of a transformer; and as primary and secondary ampere turns in a transformer are approximately equal, the current in the armature turn during commutation is very large; if not limited by the resistance or reactance of the coil, it is as many times greater than the full load current, as the field coil has turns. This causes serious sparking, if not taken care of.

One way of mitigating the effect of this short circuit current is to reduce it by interposing resistance or reactance; that is by making the leads between the armature turns and the commutator bars of high resistance or high reactance. Obviously this arrangement can merely somewhat reduce the spark-

ing by reducing the current in the short circuited coil, but can not eliminate it; and it has the disadvantage, that in the moment of starting, if the motor does not start at once, the resistance lead is liable to be burned out by excessive heating; while when running, each lead is in circuit only a very small part of the time: during the moment when the armature turn to which it connects, is under a commutator brush. As the resistance of the lead must be very much greater than that of the armature coil, and as the space available for it is very much smaller, if remaining in circuit for any length of time, it is destroyed by heat.

In direct current motors, commutation may be controlled by an interpole or commutating pole; that is, by producing a magnetic field at the brush, in direction opposite to the field of armature reaction, and by this field inducing in the armature turn during commutation, an e. m. f. of rotation which reverses the current. Such a commutating pole, connected in series into a circuit, would, in the alternating current motor, induce an e. m. f. in the short circuited turn, by its rotation; but this e. m. f. would be in phase with the field of the commutating pole, and thus with the current, that is, with the main field of the motor. Therefore it could not neutralize the e. m. f. induced in the short circuited turn by the alternation of the main field through it, since this latter e. m. f. is in quadrature with the main field, and thus with the current; but would simply add itself to it, and so make the sparking worse. A series commutating pole, while effective in a direct current motor, is therefore ineffective in an alternating current motor, due to its wrong phase.

To neutralize the e. m. f. induced by the alternation of the main field through the armature turn during commutation,

by an opposite e. m. f. induced in this turn by its rotation through a quadrature field or commutating field, this field must therefore have the proper phase. The e. m. f. of alternation of the main field through the short circuited turn is proportional to the main field F and frequency N , and is in quadrature with the main field. The e. m. f. induced in the short circuited turn by its rotation through the commutating field is proportional to the frequency of rotation or speed N_0 , and to the commutating field F_0 , and in phase therewith; to be in opposition and equal to the e. m. f. of alternation, the commutating field must therefore be in quadrature with the main field, and frequency times main field must equal speed times commutating field. That is:

$$N F = N_0 F_0$$

or in other words, the commutating field must be:

$$F_0 = \frac{N}{N_0} F$$

or equal to the main field times the ratio of frequency to speed, and in quadrature therewith.

Hence, at synchronism: $N_0 = N$, the commutating field must be equal to the main field; at half synchronism: $N_0 = \frac{1}{2} N$, it must be twice; at double synchronism: $N_0 = 2 N$, it must be one-half the main field.

The problem of controlling the commutation of the alternating current motor therefore requires the production of a commutating field of proper strength, in quadrature phase with the main field of the motor, and thus with the current.

In a transformer, on non-inductive or nearly non-inductive secondary load, the magnetism is approximately in quadrature behind the primary, and ahead of the secondary current; transformation between compensating winding and arma-

ture thus offers a means of producing a quadrature field in the alternating current motor for compensation.

In the conductively compensated series motor, at perfect compensation, no quadrature field exists; while with over or under compensation, a quadrature field exists, in phase with the current, and therefore not effective as commutating field.

In the inductively compensated series motor, the quadrature field, which transforms current from the armature to the compensating winding, is of negligible intensity, as the compensating winding is short circuited, and thus consumes very little voltage.

A quadrature field, however, appears in those motors in which the compensating winding is primary, and the armature secondary, that is in repulsion motors; since in the armature the induced or transformer e. m. f. is opposed by the e. m. f. of rotation; so a considerable e. m. f. is induced, and therefore a considerable transformer flux exists.

Therefore, when impressing the supply voltage on the compensating winding, and short circuiting the armature upon itself, that is, in the repulsion motor, the voltage is supplied to the armature by transformation from the compensating winding, and the magnetic flux of this transformer is in quadrature with the supply current; that is, it has the proper phase as commutating flux.

The repulsion motor thus has in addition to the main field, in phase with the current, a transformer field, in quadrature with the main field in space and in time, and so in the proper direction and phase as commutating field; thus giving perfect commutation if this transformer field has the intensity required for commutation, as discussed above.

As in the repulsion motor, the armature is short circuited upon itself, the voltage supplied to it by transformation from

the compensating winding equals the voltage consumed in it by the rotation through the main field. The former voltage is proportional to the frequency N and to the transformer field F^1 , the latter to the speed N_0 and to the main field F , and it so is:

$$N F^1 = N_0 F,$$

that is, the transformer field is:

$$F^1 = \frac{N_0}{N} F$$

or equal to the main field times the ratio of speed to frequency.

Comparing this value of the transformer field of the repulsion motor, F^1 , with the required commutating field F_0 , it is seen that at synchronism $N_0 = N$, $F^1 = F_0$; that is, the transformer field of the repulsion motor has the proper value as commutating field, so that no short circuit current is produced in the armature turn under the brush, but the commutation is as good as in a direct current motor with negligible armature reaction.

At half synchronism, $N_0 = \frac{1}{2} N$, the transformer field of the repulsion motor: $F^1 = \frac{1}{2} F$, is only one quarter as large as the commutating field required $F_0 = 2 F$, and the short circuit current is reduced by 25% below the value which it has in the series motor; and the commutation, while it is better, is not yet perfect.

At double synchronism: $N_0 = 2 N$, the transformer field is $F^1 = 2 F$, while the commutation field should be: $F_0 = \frac{1}{2} F$, and the transformer field thus is four times larger than it should be for commutation; so that only one-quarter of the transformer field is used to neutralize the e. m. f. of alternation in the short circuited turn; the other three-quarters induces an e. m. f., thus causing a short circuit current three times as large as it would be in a series motor. That is, the short circuit current under the brush, and thus the sparking, in

the repulsion motor at double synchronism is very much worse than in the series motor, and the repulsion motor at these high speeds is practically inoperative.

Hence, as regards commutation, a repulsion motor is equal to the series motor at standstill where no compensation of the short circuit current is possible—but becomes better with increasing speed: as good as a good direct current series motor at synchronism; and then again becomes worse by over compensation, until at some speed, at 40% above synchronism, it again becomes as poor as the alternating current series motor; above this speed, it becomes rapidly inferior to the series motor.

To produce right intensity of the transformer field, to act as commutating field, it is therefore necessary above synchronism to reduce the transformer field below the value which it would have when transforming the total supply voltage from compensating winding to armature. This means, that above synchronism, only a part of the supply voltage must be transformed from compensating winding to armature, the rest directly impressed upon the armature. Thus at double synchronism, where the transformer field of the repulsion motor is four times as strong as is required for commutation, to reduce it to one-quarter, only one-quarter of the supply voltage must be impressed upon the compensating winding, three-quarters directly on the armature.

To get zero short circuit current in the armature turn under the brush, below synchronism more than the full supply voltage would have to be impressed upon the compensating winding, which usually cannot conveniently be done. At synchronism the full supply voltage is impressed upon the compensating winding, while the armature is short circuited as repulsion motor; and with increasing speed above synchronism,

more and more of the supply voltage is shifted over from compensating winding to armature; that is, the voltage impressed upon the compensating winding is reduced, from full voltage at synchronism, while the voltage impressed upon the armature is increased, from zero at synchronism, to about three-quarters of the supply voltage at double synchronism. Such a motor, in which the transformer field is varied in accordance with the requirement of commutation, is called a "series repulsion motor."

The arrangement described here eliminates the short circuit current induced in the commutated armature turn by the alternation of the main field, and that completely above synchronism, so that during commutation, no current is induced in the armature turn. This, however, is not sufficient for perfect commutation: during the passage of the armature turn under the brush, the current in the turn should reverse; so that in the moment in which the turn leaves the brush, the current has already reversed. For sparkless commutation, it therefore is necessary, in addition to the neutralizing e. m. f. of the transformer field, to induce an e. m. f. which reverses the current. This e. m. f., and thus the magnetic flux which induces it by the rotation, must be in phase with the current. That is, in addition to the "neutralizing" component of the commutating field (which is in quadrature with the current), to reverse the current, a second component of the commutating field must exist, in phase with the current; this component so may be called the "reversing field". The total commutating field required to eliminate the short circuit current due to the alternating main field by the "neutralizing" flux, and to reverse the armature current by the "reversing flux", must therefore be somewhat less than 90° lagging behind the main field and thus the main current.

While in a transformer with non-inductive load on the secondary, the magnetic flux lags nearly 90° behind the primary current, in a transformer with inductive load on the secondary, the magnetic flux lags less than 90° behind the primary current; and the more so the higher the inductivity of the secondary load.

Therefore, by putting a reactance into the armature circuit of the motor, and so making the armature circuit inductive, the transformer flux is made to lag less than 90° behind the current, and act not only as neutralizing but also as reversing flux; and so, if it be of proper intensity, it gives perfect commutation.

An additional reactance would in general be objectional, in lowering the power factor of the motor. The motor, however, contains a reactance: its field circuit, which has to be excited, can be used as reactance for the armature circuit. That is, by connecting the field coils into the armature circuit, or in other words, using secondary excitation, the transformer flux of the motor is given the lead ahead of quadrature position with the main field, which is required to act as reversing field.

In this manner, it is possible in the alternating current commutator motor, to get at all speeds from synchronism upwards, the same perfect commutation as in a direct current motor with commutating poles, by varying the distribution of supply voltage between compensating winding and armature, and exciting the field in series with the armature circuit; that is, in the series repulsion motor B of the preceding table. Obviously, this distribution of voltage would for all practical purposes be carried out sufficiently by using a number of steps, as shown diagrammatically by the arrangement in Fig. 46:

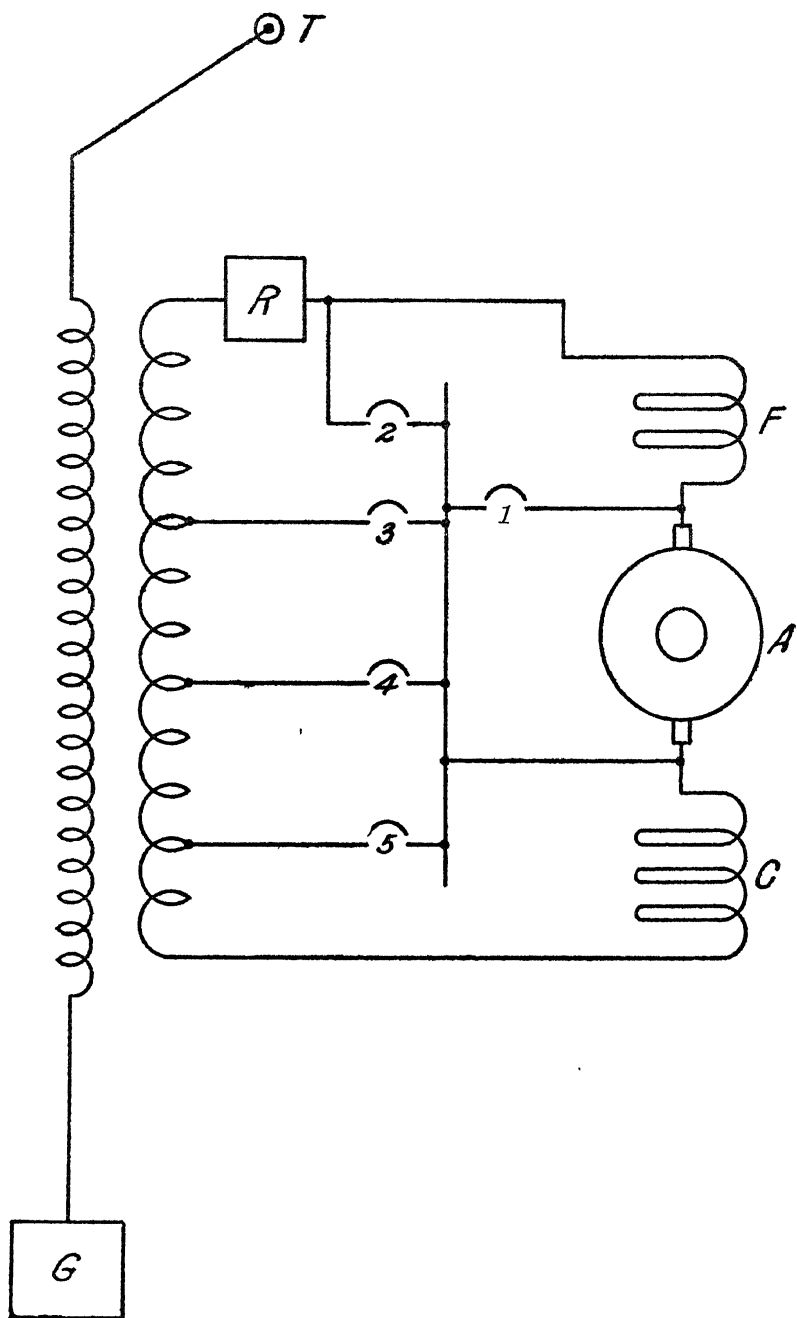


Fig. 46.

T is the supply circuit, F the field winding, A the armature, and C the compensating winding.

Closing switch 1, and leaving all others open, the motor is a repulsion motor.

Closing switch 2, and leaving 1 and all others open, the motor is a repulsion motor with secondary excitation.

Closing switch 3, or 4, 5 . . . and leaving all others open, the motor is a series repulsion motor B, with gradually increasing armature voltage and decreasing voltage on the compensating winding.

By winding the armature for half the voltage and twice the current of the compensating winding, when changing from position 1, the field in the compensating circuit, to the next position, with the field in the armature circuit, the field current and the field strength becomes double the value it had in starting, where no compensation exists, and which it would have to maintain in a series motor; and thus a correspondingly greater motor output is secured, than would be possible in a motor in which the commutation is not controlled.

FIFTEENTH LECTURE

ELECTROCHEMISTRY

ELECTROCHEMISTRY is one of the most important applications of electric power, and possibly even more power is used for electrochemical work than for rail-roading.

In electrochemical industries the most expensive part is electric power; material and labor are usually much less. Such industries therefore are located at water powers, where the cost of power is very low.

The main classes of electrochemical work are :

A. Electrolytic.

B. Electrometallurgical.

A. ELECTROLYTIC WORK.

The chemical action of the current is used, by electrolyzing either solutions of salts or fused salts or compounds.

Electrolysis of solutions in water is possible only with such metals which have less chemical affinity than hydrogen. For instance, Cu, Fe, and Zn can be deposited from salt solutions in water, but not Al, Mg, Na, etc. Electrolyzing, for instance, NaCl (salt solution) the sodium (Na) which appears at the negative terminal immediately dissociates the water and gives $\text{Na} + \text{H}_2\text{O} = \text{NaOH} + \text{H}$, or: sodium plus water = caustic soda plus hydrogen.

It takes 1.4 volts to electrolyze water; any metal requiring more than 1.4 volts for separation therefore is not separated, but hydrogen is produced.

Therefore the highest voltage used in an electrolytic cell containing water is 1.4 + the *ir* drop in the resistance of the

cell; which latter, for reasons of economy, is made as low as possible.

Even fused salts require fairly low voltage, at the highest from 3 to 4 volts.

Since the voltage required per cell is very low, a large number of cells are connected in series, and even then large low voltage machines are required.

Some of the important applications of electrolysis are :

Electroplating; that is, covering with copper, nickel, silver, gold, etc.

Electrotyping; that is, making of copies, usually of copper; and especially

Metal refining.

A very large part of all the copper used is electrically refined. The crude copper as cast plate is used as anode or positive, and a thin plate of refined copper is used as cathode, or negative terminal in a copper sulphate solution. The anode is dissolved by the current and the fine copper is deposited on the cathode; while silver and gold go down into the mud, lead goes into the mud as sulphate, tin as oxide; sulphur, selenium and tellurium, arsenic and other impurities also go in the mud; and zinc and iron remain in solution as sulphates if the current density is sufficiently low. If the current density is high, some zinc and iron may deposit: zinc and iron have a greater chemical activity than copper, since they precipitate copper from solution. Therefore it takes more power, that is, more voltage, to deposit zinc and iron, than it takes to deposit copper. If the current density is low, the voltage required to deposit the copper plus the *ir* drop, that is, the total voltage of the cell, is less than the voltage required to deposit zinc or iron, and they do not deposit, but ~~dissolve~~ at the anode and remain in solution.

At higher current density the *ir* drop in the cell is higher; thus the total voltage of the cell is higher, and may become high enough to deposit iron or even zinc.

If the anode is crude copper, the cathode pure copper, the voltage at the anode is higher than at the cathode and the cell takes some voltage. The voltage required for copper refining is the higher, the more impure the copper is; but is always very low, usually a fraction of a volt, and therefore very many cells are run in series.

The solution gradually becomes impure and has to be replaced.

Other metals are occasionally refined electrolytically, but only to a small extent.

Metal Reduction.

Metals are reduced from their ores electrolytically, especially such metals which have so high chemical affinity that they are not reduced by heating with carbon. In this way aluminum, magnesium, sodium, calcium, etc., are made electrolytically. Since their chemical affinity is greater than that of hydrogen, they cannot be deposited from solutions in water, but only from fused salts, or solutions in fused salts. So calcium is produced now by electrolyzing fused calcium chloride, CaCl_2 . Aluminum is made by electrolyzing a solution of alumina in melted cryolite (sodium aluminum fluoride).

SECONDARY PRODUCTS.

Frequently electrolysis is used to produce not the substances which are directly deposited, but substances produced by the reaction of these deposits on the solutions. For instance, electrolyzing a solution of salt, NaCl , in water, we get sodium, Na , at the negative, chlorine, Cl , at the positive terminal.

If we use mercury, Hg, as negative electrode, it dissolves the sodium and so we get sodium amalgam.

Otherwise the sodium does not deposit but immediately acts upon the water and forms sodium hydrate or caustic soda, NaOH.

The chlorine, Cl, at the anode also reacts on the water, one chlorine atom taking up one hydrogen and another chlorine atom the remaining OH of the water H₂O; that is, we get $2\text{Cl} + \text{H}_2\text{O} = \text{ClH} + \text{ClOH}$, that is, hydrochloric + hypochlorous acid.

With the sodium hydrate from the other cathode these acids form NaCl and ClONa, that is sodium chloride and hypochlorite, or bleaching soda.

If the solution is hot, the reaction goes further and we get $6\text{Cl} + 3\text{H}_2\text{O} = 5\text{ClH} + \text{ClO}_3\text{H}$, that is hydrochloric and chloric acid, and with the sodium hydrate from the other side these form NaCl and ClO₃Na, that is, sodium chloride and sodium chlorate.

In this way considerable industries have developed, producing electrolytically caustic soda, bleaching soda, and chlorates.

Alternating current is used very little for electrolytic work, as with organic compounds to produce oxidation and reduction at the same time; that is, act on the compound in rapid succession by oxygen and hydrogen, the one during the one, the other during the next half wave of current.

Very active metals like manganese and silicon dissolve by alternating current; that is, one-half wave dissolves, but the other does not deposit again.

Very inert metals like platinum are deposited by alternating current; that is, the negative half wave deposits by alternating current, but the positive half wave does not dissolve.

B. ELECTROMETALLURGICAL WORK.

In electrometallurgical work the heat is used to produce the chemical action; thus it is immaterial whether alternating or direct current is used.

The voltage required is still low but not as low as in electrolytic work:

The carborundum furnace takes from 250 to 90, mostly about 100 volts; that is, it starts cold with 250 volts. While heating up the resistance drops, and the voltage decreases down to 100 volts when the furnace is hot and remains there until towards the end. Then the inner layer of carborundum begins to change to graphite and the resistance, and therefore the voltage falls.

The carbide furnace and arc furnaces in general take from 50 to 100 volts; the graphite furnace takes from 10 to 20 volts.

To get very high temperatures a very large amount of energy has to be concentrated in one furnace; and with the moderate voltage used, this requires very large currents, thousands of amperes. Alternating currents are almost exclusively used, since it is easier to produce very large alternating currents by transformers, and since it is easier to control alternating than direct currents.

Electric heat necessarily is very much more expensive than heat produced by burning coal, and so the electric furnace is used mainly:

1st. Where very perfect control of the temperatures and freedom from impurities is essential.

2nd. Where temperatures higher than can be produced by combustion are required.

1. Very accurate temperature regulation and freedom from impurities, for instance, are important in making and

annealing high grade tool steels, etc. By using coal or oil as fuel, contamination by the gases of combustion, and by the metal taking up carbon or (if an excess of air is used, oxygen) is difficult to avoid.

By electric heating, by resistance at lower temperature and by induction furnace at higher temperature, contamination can be perfectly avoided and even the air can be excluded.

2. The temperature of combustion is limited.

Four-fifths of air is nitrogen which does not take part in the combustion, but which has to be heated, thus greatly lowering the temperature; therefore combustion in air, even if the air is preheated, gives a lower temperature than when using oxygen. But even the temperature of the oxy-hydrogen, or the oxy-acetylene flame is only just able to melt platinum.

The temperature which can be reached by combustion, is limited, since at very high temperature the chemical affinity of oxygen for hydrogen and carbon ceases: water dissociates, that is, spontaneously splits up in hydrogen and oxygen at 2000 degrees Centigrade and no temperature higher than 2000° can therefore be reached by the oxy-hydrogen flame; carbon dioxide, CO_2 , already dissociates at about 1500°C into carbon monoxide, CO , and oxygen, O . Carbon monoxide, CO , splits up into carbon and oxygen not much above 2000°C. (In all high temperature reactions of carbon, as in the formation of carbides, CO therefore always forms and not CO_2 , since CO_2 cannot exist at a very high temperature; and the CO when leaving the furnace then burns to CO_2 with blue flame).

Higher temperatures than those generated by the combustion of carbon and hydrogen can be produced by the combustion of those elements whose oxides are stable at very high temperatures, as aluminum and calcium. In this way, many metals, as chromium and manganese, which cannot be reduced

from the oxides by carbon (due to the lower temperature of carbon combustion) can be reduced by aluminum in the "thermite" process. That is, their oxides are mixed with powdered aluminum and then ignited: the aluminum burns in taking up the oxygen of the metal, and so produces an extremely high temperature, which melts the metal and the alumina (corundum) which is produced.

Since, however, all the aluminum is made electrolytically, the thermite process still requires the use of electric power. The temperature of combustion of aluminum, however, is still far below that of the electric carbon arc, since in the carbon arc, alumina boils.

For temperatures above 2000° to 2500°C , and up to the arc temperature or about 3500°C , electric energy is therefore necessary.

Electric furnaces are of two classes:

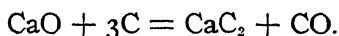
Arc Furnaces and Resistance Furnaces.

In the resistance furnace any temperature can be produced up to the point of destruction of the resistance material, that is, up to 3500°C , when using carbon.

The arc furnace gives the arc temperature of 3500°C , but allows the concentration of much more energy in a small space and thus produces reactions requiring the very highest temperatures.

Some of the electrometallurgical industries are:

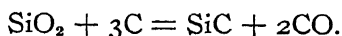
(a). Calcium carbide production. Arc furnaces are used and the reaction is



A mixture of coke and quick lime is used in the process.

(b). Carborundum production. A resistance furnace is used, containing a carbon core about 24 feet long, around which the material is placed and heated by the current passing

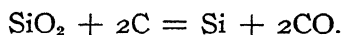
through the core. The furnace takes 1000 HP and the reaction is:



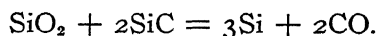
The material is a mixture of sand, coke, sawdust and salt.

(c). Graphite furnace. A resistance furnace somewhat similar to the carborundum furnace is used, but with lower voltage and larger currents; the material is coke or anthracite, which by the high temperature is converted into graphite, probably passing through an intermediate stage as a metal carbide.

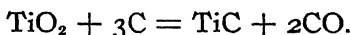
(d). Silicon furnace. Either arc or resistance furnace is used; the reaction is:



or,



(e). Titanium carbide furnace. Arc or resistance furnace is used which requires a very high temperature; that is, a greater temperature than that of the calcium carbide furnace.



Other products of the electric furnace are siloxicon, silicon monoxide, etc., and numerous alloys of refractory metals, mainly with iron; as of vanadium, tungsten, molybdenum, titanium, etc., which are used in steel manufacture.

The use of the electric arc for the production of nitric acid and nitrate fertilizers; of the high potential glow discharge for the production of ozone for water purification, etc., also are applications of electric power, which are of rapidly increasing industrial importance.

SIXTEENTH LECTURE

THE INCANDESCENT LAMP

THE two main types of electric illuminants are the incandescent lamp and the arc.

In the incandescent lamp the current flows through a solid conductor, usually in a vacuum, and the heat produced in the resistance of the conductor makes it incandescent, thus giving the light. Incandescent lamps in an electric circuit therefore act as non-inductive ohmic resistance and can therefore be operated equally well on constant potential as on constant current. As electric distribution systems are always constant potential, most incandescent lamps are operated on constant potential; and only for outdoor lighting, that is, for street lighting in cases where the arc lamp is too large and too expensive a unit of light for the requirements, incandescent lamps are used on a constant, direct or alternating current circuit; they are then usually built for the standard arc circuits, and thus for low voltage.

For general convenience the efficiency of incandescent lamps is given in watts power consumption per horizontal candle power, when operating on such a voltage, that the candle power of the lamp decreases by 20% in 500 hours running; and the time, in which the candle power decreases by 20%—that is, 500 hours with the present efficiency rating—is called the useful life; since experience has shown, that after a decrease of candle power of 20%, with the carbon filament lamp, under average conditions, it is more economical to replace the lamp with a new lamp, than to continue its use; as then the increased cost of light due to the lower efficiency is greater than the cost of the lamp, when distributed over 500 hours.

In discussing incandescent lamp efficiencies, it is therefore essential to make sure that the efficiency is given at the useful life of 500 hours; since obviously any efficiency can be produced in any lamp, by running it at higher voltage, but the life is greatly shortened thereby. Therefore efficiency comparisons have a meaning only when based on the same length of useful life, as 500 hours.

Obviously, for other types of lamps, the economic life may be greater (as for more expensive lamps) or less than 500 hours.

Illuminants are measured and compared by the total flux of light which they give. Usually, however, this is expressed in "mean spherical candle power"; that is, the candle power which would be given by the illuminant if this light were distributed uniformly throughout. Since the object of a lamp is to give light, obviously the only logical way of measuring it is by the total amount of light which it gives, and so by the mean spherical candle power; this therefore is standard.

The conventional rating of the incandescent lamp, in horizontal candle power, therefore has to be multiplied by a reduction factor, to give the mean spherical candle power. With the carbon filament lamp, this reduction factor is usually .79; a 16 candle power so has a mean spherical candle power of $16 \times .79 = 12.6$ c. p., and at an efficiency of 3.1 watts per horizontal candle power, it has an efficiency of $\frac{3.1}{.79} = 3.92$ watts per mean spherical candle power.

The carbonized bamboo fibre used in the very early days was very soon replaced by filaments made of structureless cellulose, squirted from a cellulose solution, and then carbonized. By "treating" these filaments, that is, heating them in gasolene vapor and therefrom depositing a thin shell of car-

bon on them, a considerable increase in efficiency became possible; their efficiency was thus greatly increased, from 5 to 6 watts per candle power in the early days, to 3.5 and 3.1 watts per candle power. Of these two types, the 3.5 watt lamp is used in systems of poor voltage regulation, in which the more efficient 3.1 watt lamp would have too short a life; with the improvements in the voltage regulation of systems, the less efficient 3.5 watt lamp is thus going out of use.

By exposing these "treated" filaments to the highest temperature of the electric furnace, their stability at high temperature is greatly improved; so that in these "metallized"* filament lamps an efficiency of 2.5 to 2.6 watts per candle power is reached. Whether a still further increase of efficiency of the carbon filament will occur, as is quite possible, or whether the carbon filament will be replaced by the metal filaments, remains for the future to decide.

In the last years, metal filament lamps giving efficiencies far higher than has so far been possible to reach with the carbon filament, have been developed. First came the osmium lamp, of 1.5 watts per candle power. As the total supply of osmium available on the earth is far less than would be required for one year's production of incandescent lamps, the osmium lamp never could hope for more than a very limited use. The tantalum lamp, which was developed next, and is now quite extensively used, gives an efficiency of about 2 watts per candle power; that is, it is not quite as efficient as the osmium lamp, since tantalum is somewhat more fusible than osmium. As tantalum is a metal which can be drawn into wire, the tantalum filament is of drawn wire; while

* The name "metallized" is given to the form of carbon produced in these filaments by the electric furnace temperature, since it has metallic resistance characteristics; a positive temperature coefficient of resistance, while the other forms of carbon have a negative temperature coefficient.

all the other metals which are used for lamp filaments are not ductile, and the filaments have to be made by some squirting process, similar to the carbon filament. The highest efficiency was reached by the tungsten (wolfram) lamp, of 1 to $1\frac{1}{4}$ watts per candle power; that is, tungsten (or rather wolfram metal, since tungsten is the name of the ore of the metal), has the highest melting point of all known metals, and so can be run at the highest temperature, that is, highest efficiency.

All these metals melt far below the temperature where carbon melts or boils, but carbon has the great disadvantage of evaporating considerably below its melting point, while these metals evaporate very little, and so can be run at a temperature fairly close to their melting point; while the carbon filament has to be operated at a temperature very far below the melting point.

The great difficulty with all these metal filaments is, that the metals are very much better conductors than carbon; to get the same filament resistance, so as to consume the same current, at the same voltage, the metal filaments must be very much longer and very much thinner than the carbon filament. As the efficiency of the metal filament is far higher, to produce the same candle power at the same voltage, less current and therefore a higher resistance is required, which makes these metal filaments still thinner; as a result, although the metals are mechanically stronger than carbon, the metal filaments are far more frail, due to their exceeding thinness, and it is very difficult to produce lamps of as low candle power, as is feasible with carbon filaments. For larger units, however, and for larger current low voltage lamps, for series lighting, the metal filaments are specially suited.

For general use, the 16 candle power lamp has proved the most convenient unit of light. The limitation of voltage, for

which efficient incandescent lamps of such size can be built, has been the cause of the general use of 110 volt distribution. 220 volt 16 candle power carbon filament lamps can be built, but are of necessity less efficient, by about 15%, than 110 volt lamps: at 220 volts, half the current and so four times the resistance is required for the same power as at 110 volts; the filament therefore is about twice as long and half as thick, hence more breakable and more rapidly disintegrating; so that there is no possibility of reaching the same efficiency in a 220 volt 16 candle power lamp, as in 110 volt lamp made with the same care. For the same reason, the 8 candle power 110 volt lamp must be less efficient than the 16 candle power 110 volt lamp.

In an incandescent lamp are specified: the candle power, the efficiency, and the voltage. To produce lamps fulfilling simultaneously all three conditions, requires either to allow a large margin in either condition—that is, gives a product inferior in uniformity—or to get a uniform product, a large percentage is thrown out as defective, and the cost of the lamp is thus seriously increased. For this reason, in the manufacture a very close agreement is aimed at in candle power and in efficiency; the lamps are then assorted for voltages, and different voltages are then assigned by the organization of illuminating companies to the different companies, so as to consume the total lamp product. As a result hereof, a far more uniform product is derived than could be derived in any other way, and than is available in any other country. This is the reason, that in distribution systems not one and the same voltage, as 110, is employed throughout; but different cities use different voltages, between 105 and 130. The average incandescent lamp used in this country therefore is decidedly superior in uniformity and in efficiency to those used abroad. The ultimate cause hereof

is, that since the earliest days the illuminating companies have followed the principle of supplying light, and not power;* and 220 volt distribution, while being more efficient from the generating station to the customer's meter, is decidedly inferior in efficiency from the generating station to the candle power produced at the customer's lamps, as the saving in distribution losses does not make up for the lower efficiency of the 220 volt lamp. For this reason, 220 volt distribution has never found any entrance in this country.

In gas lighting, an enormous increase of efficiency resulted from the development of the Welsbach gas mantle. In the same direction, that is, by using what may be called "heat luminescence" in electric lighting, the Nernst lamp was developed, using the same class of material: refractory metallic oxides, as in the Welsbach mantle. The "glower" of the Nernst lamp, however, is a non-conductor at ordinary temperature, and requires some heating device, the "heater", to be made conducting. When conducting, it has a very high negative temperature coefficient; that is, the voltage consumed by the glower decreases with the increase of current, just as in the arc, and it therefore requires a steadying resistance, called the "ballast". The lamp therefore requires some operating mechanism, to cut the heater out of circuit after the glower is started. The glower of the Nernst lamp is not operative in a vacuum, since air seems to be necessary for its heat luminescence. Fairly good efficiencies have been reached with these lamps, especially in larger units, as 3 to 6 glower lamps, but not of the same class as with the tungsten lamp.

* For a long time, the bills were even made out in "lamp hours" and in the earlier days the machines rated in "lights" and not in kilowatts.

SEVENTEENTH LECTURE

ARC LIGHTING

WHILE incandescent lamps can be operated on constant potential as well as on constant current, the arc is essentially a constant current phenomenon. At constant length, the voltage consumed by the arc decreases with increase of current, as shown by curve I in Fig. 47. If, therefore, an attempt is made to operate such an arc on constant potential, for instance on 80 volts—which would correspond to 3.9 amperes on curve I—then any tendency of the current to increase—as by a momentary drop of the arc resistance—would lower the required arc voltage, and so increase the current, at constant supply voltage, hence still further lower the arc voltage, etc., and a short circuit would result. Vice versa, a momentary decrease of arc current, by requiring more voltage than is available, would still further decrease the current, increase the required voltage, etc., and the arc would extinguish.

Therefore only such apparatus is operative on constant potential, in which an increase of current requires an increase of voltage, and vice versa; and so limits itself.

While therefore arcs can be operated on a constant current system, to run arc lamps on constant potential, some current limiting device is necessary in series with the arc, as a resistance; or, in an alternating current circuit, a reactance.

The voltage consumed by the resistance is proportional to the current, and a resistance of 8 ohms inserted in series to the arc would thus consume the voltage shown in straight line II in Fig. 47. The voltage consumed by the arc plus the resistance then is given by the curve III, derived by adding I and II. As

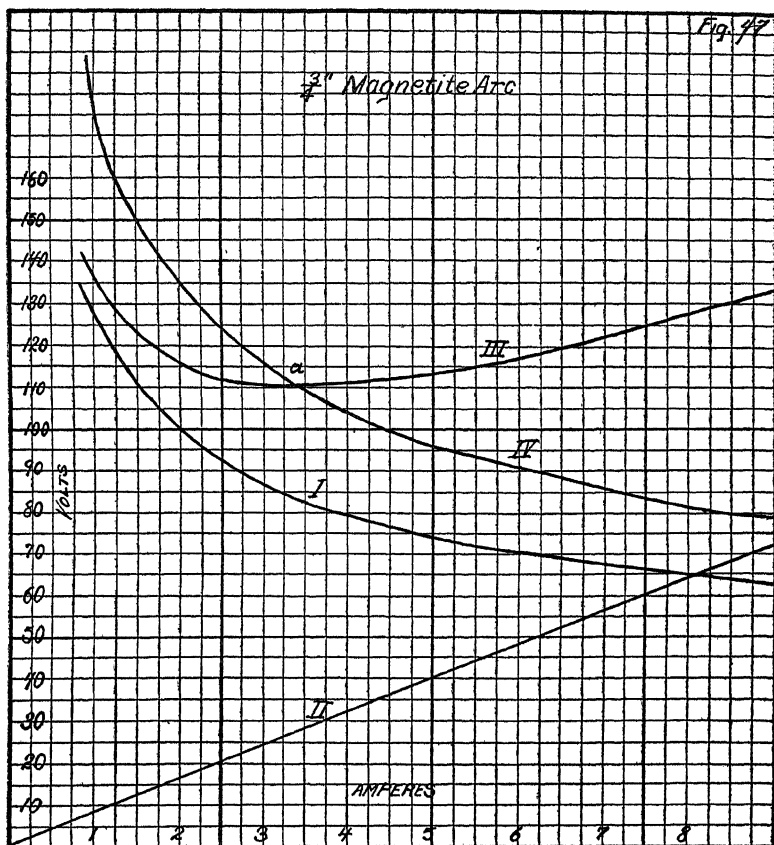


Fig. 47.

seen, below 3.35 amperes, the total required voltage still decreases with increase of current, and the arc is still unstable; that is, the resistance is insufficient. Above this current, an increase of current requires an increase of voltage and so limits itself; that is, the arc is stable; with 8 ohms series resistance, 3.35 amperes therefore is the limit of stability of the arc; and attempting to operate it at lower current, as for instance at 2 amperes and 116 volts supply, the arc either goes out, or

the current runs up to 5.5 amperes, where the arc becomes stable on 116 volts supply.

With a higher series resistance, the arc remains stable to lower currents, and vice versa. It follows herefrom, that for the operation of an arc lamp on constant potential, a higher voltage is required than that consumed by the arc proper.

At every value of series resistance therefore a point *a* in Fig. 47, is reached, at which for decreasing current the arc becomes unstable; and all these points, for different resistance values, give a curve IV, which is called the "stability curve" of the arc curve I.

The supply voltage required to operate the arc represented by curve I must therefore be higher than that given by the stability curve IV. For instance, at 4 amperes, the arc cannot be operated at less than 104 volts supply. At 104 volts supply the limit of stability is reached; that is, a change of current does not require a change of voltage, but the arc voltage decreases as much as the resistance voltage increases and the current thus drifts; and for supply voltages higher than 104, the arc is stable, the more so, the higher the supply voltage is above 104. The difference in voltage between the supply voltage and the arc voltage thus is consumed by the "steading resistance" of the arc.

High reactance in series with the direct current arc retards the current fluctuations and so reduces them; so that with reactance in series to the direct current arc, the arc can be operated by a supply voltage closer to the stability curve IV than without reactance; reactance therefore is very essential in the steadying resistance of a direct current arc. Obviously, no series reactance can enable operation of the arc I on a supply voltage below that given by the stability curve IV.

The arc characteristic I is far steeper for low currents than for high currents, and is the steeper the greater the arc length. Low current arcs and long arcs therefore require, that on a constant potential supply, a greater part of the supply voltage is consumed by the steadying resistance (or steadying reactance with alternating arcs) than high current arcs, or short arcs; and are therefore less economical on constant potential supply.

Constant potential arc lamps are necessarily less efficient than constant current arc lamps, due to the power consumed in the steadying resistance. A large part of this power is saved in alternating constant potential arc lamps, by using reactance instead of resistance, but the power factor is therefore greatly lowered; that is, the constant potential alternating arc lamp rarely has a power factor of over 70%.

Where therefore high potential constant current circuits are permissible, as for outdoor or street lighting, arc lamps are usually operated on a constant current circuit, with series connection of from 50 to 100 lamps on one circuit. With the exception of a few of the larger cities, all the street lighting by arc lamps in this country is done by constant current systems, either direct current or alternating current.

For direct current constant current supply, separate arc light machines have been built, and are still largely used. In these machines, inherent regulation for constant current is produced by using a very high armature reaction and relatively weak field excitation; that is, the armature ampere turns are nearly equal and opposite to the field ampere turns, and thus both very large compared with the difference, the resultant ampere turns, which produce the magnetic field. A moderate increase of current and consequent increase of armature ampere turns therefore greatly reduces the resultant ampere turns and

so the field magnetism and the voltage, that is, the machine tends to regulate for constant current. Perfect constant current regulation then is secured by some governing device, as an automatic regulator varying a resistance shunted across the series field. It must, however, be understood that the "regulator" of the arc machine does not give a constant current regulation, but the armature reaction of the machine does this, and the regulator merely makes it perfect; but even with the regulator disconnected, arc machines give fairly close constant current regulation.

As the voltages produced by arc machines are very high—4,000 to 10,000—commutation of the current, with the ordinary commutator, which is limited to a maximum of 40 to 50 volts per segment—is not well suited, but rectification is used. The Brush arc machine therefore is a quarter-phase alternator with rectifying commutator. That is, the commutator shifts the connection over from the phase of falling e. m. f. to that of rising e. m. f., and thereby is able to control as high as 3,000 volts per commutator ring.

With the development of the mercury arc rectifier, which converts constant alternating current into constant direct current, arc machines are going out of use. The arc machine necessarily must be a small unit, since 100 to 150 lamps in series give already as high a voltage as is safe to use in arc circuits, but do not yet represent much power; and when supplying thousands of arc lamps a large number of small machine units are required, which are uneconomical in space, in attendance and in efficiency. The mercury arc rectifier in combination with the stationary constant current transformer enables us to derive the power from the alternating current constant potential supply system.

Constant alternating current is derived by a constant current transformer or constant current reactance. Diagrammatically, the constant current transformer is shown in Fig. 48. The primary coil P and the secondary coil S are movable

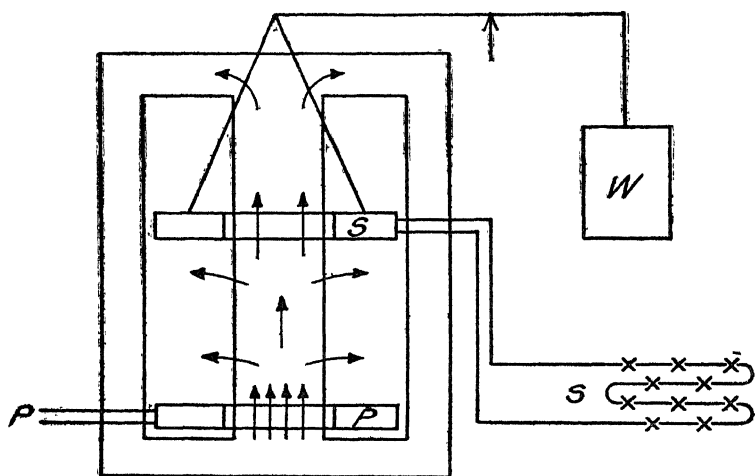


Fig. 48.

with regard to each other (which of the two coils is movable, is immaterial, or rather, is determined by consideration of design). Fig. 48 shows the coil S suspended and its weight partially balanced by counter-weight W.

With the secondary coil S close to the coil P, that is, in the lowest position, most of the magnetism produced by the primary coil P passes through the secondary coil S, and the secondary voltage therefore is a maximum. The further the secondary coil moves away from the primary coil, the more of the magnetism passes between the coils, the less through the secondary coil, and the lower therefore is the secondary voltage, which becomes a minimum (or zero, if so desired), with the

secondary coil at a maximum distance from the primary, that is, in the top position.

Primary current and secondary current are proportional and in opposition to each other, and repel each other, and the repulsion is proportional to the product of the two currents; that is, proportional to the square of the secondary current. The weight of the secondary coil is balanced by the counter-weight *W* and the repulsion from the primary coil, at normal secondary current. Any increase of secondary current by a decrease of load, increases the repulsion, in this way pushing the secondary coil further away from the primary and thereby reducing the secondary voltage and thus the current; and vice versa, a decrease of secondary current, by an increase of load, reduces the repulsion and so causes the secondary coil to come nearer to the primary, that is, increases its voltage and so restores the current. Such an arrangement regulates for constant current between the voltage limits given by the two extreme positions of the movable coil. These usually are chosen from some margin above full load, down to about one-third load.

The constant current reactance operates on the same principle: the two coils *P* and *S* are connected in series with each other into the arc circuit supplied from the constant potential source, and by separating or coming together, vary in reactance with the load, and thereby maintain constant current.

While the alternating current arc lamp is less efficient, that is, gives less light for the same power, than the direct current arc lamp, the disadvantages of the use of numerous arc machines have led to the extended adoption of alternating current series arc lighting before the development of the mercury

arc rectifier, which enabled the operation of direct current arc circuits from constant current transformers.

While incandescent lamps give the same efficiency for all sizes except such small sizes where mechanical difficulties appear in the filament production, the efficiency of the arc decreases greatly with decrease of current; that is, the arc is at the greatest efficiency only for large units of light, but rather inefficient and not so well suited for small units of light.

Even in large units, the efficiency of light production of the direct current carbon arc lamp is not superior to that of the tungsten incandescent lamp; that of the alternating current carbon arc lamp is inferior to the tungsten lamp; and the carbon arc lamp thus finds its field mainly where large units of light are required, especially as long as the cost of renewal of the metal filament lamps is still very great. Entirely different, however, are the conditions developed in the last years, with the luminous arcs, as the flame carbon arc, the mercury lamp, and the magnetite and titanium carbide arc. In these, efficiencies of light production have been reached which no incandescent lamp can hope to approach.

In the carbon arc, practically all the light comes from the incandescent tips of the carbons, very little from the arc flame. Then by using materials, which in the arc flame give an intensely luminous spectrum, the efficiency of the arc lamp has been vastly improved.

So far only three materials have been found, which in luminous arcs give efficiencies vastly superior to incandescence: mercury, calcium (lime), and titanium. All three even in moderate sized units, give efficiencies of one-half watt or better per candle power.

The mercury arc has the advantage of perfect steadiness, a long life—requiring no attention for thousands of hours—

and high efficiency over a fairly wide range of candle powers; but it is seriously handicapped for many purposes by its bluish-green color.

In the flame carbon lamp carbons impregnated with calcium compounds, usually calcium fluoride (fluorspar) are used, and the arc then has an orange-yellow color. Other compounds which give red or white color to the arc are so much inferior in efficiency that they are used only to a very limited extent. The compounds, after coloring the arc and giving it efficiency, escape as smoke; the arc therefore must be an open arc. This, however, means short life of the carbons and frequent trimming.

The open arc lamp, which was used formerly, has, however, been almost entirely superseded by the enclosed carbon arc, in spite of the somewhat lower efficiency of the latter; and the inconvenience of daily attendance required by an open arc, and the large consumption of carbons, makes a return to this type improbable. For this reason the flame carbon lamp has not proven suitable for general outdoor illumination, as street lighting, where the cost of carbons and trimming would usually far more than offset the gain in efficiency. Flame carbon lamps, however, have found a field for decorative lighting, for advertising purposes, etc., for which the glare of light and its color makes them very suitable. They are generally used on constant potential circuits with two or three lamps in series.

To eliminate the objections of short life and consequent frequent trimming and high cost of carbons, and thereby make the luminous arc able to enter the field of general outdoor illumination, carbon had to be eliminated altogether as electrode material, and its place was taken by magnetite, while titanium compounds give the high efficiency. This led to the long

burning luminous arc of the white color of the titanium-iron spectrum as represented by the magnetite arc, the metallic oxide arc, and other types still in development.

In all these long burning luminous arcs, some efficiency had to be sacrificed in developing sufficiently small units for general illumination. While the substitution of the flame carbon in the open arc has quadrupled the light at the same power consumption, and the substitution of the magnetite electrode for carbon at the same power consumption would in the same manner increase the light, for street illumination the main problem was, to decrease the power consumption rather than increase the amount of light given; and so in the long burning luminous arcs, which are now beginning to replace the carbon arcs of old, the power consumption has been reduced by from 30 to 60% with a sufficient increase of light to be marked.

In the arc lamp, the current is carried across the gap between the terminals by a stream of vapor of the electrodes; thus the electrodes consume more or less rapidly. Some feeding mechanism is therefore required to move the electrodes towards each other during their consumption. This arc lamp mechanism may be operated by the current, or by the voltage, or by both; this gives the three different types of lamps: the series lamp, the shunt lamp, and the differential lamp.

In the series lamp, an electromagnet energized by the lamp current, and balanced against a weight or a spring, moves the carbons towards each other when by their burning off, the arc lengthens and the current decreases. Obviously, this lamp cannot be used on constant current circuits, or with several lamps in series, but only as single lamp on constant potential circuits, and therefore has practically disappeared.

In the shunt lamp, the controlling magnet is shunted across the arc, and with increasing arc length and consequent

arc voltage, moves the electrodes towards each other. In constant current circuits, this lamp tends towards hunting, and therefore requires a very high reactance in series; it thereby gives a lower power factor in alternating current circuits, and has therefore been superseded by the differential lamp. It has, however, the advantage of not being sensitive to changes of current.

In the differential lamp, an electromagnet in series with the arc opposes an electromagnet in shunt to the arc, and the lamp regulates for constant arc resistance. It is the lamp now universally used in constant potential and constant current systems, is most stable in its operation; but in constant current systems, it requires that the current be constant within close limits: if the current is low, the arc is too short, and the lamp gives very little light, and if the current is high, the arc becomes so long as to endanger the lamp.

From the operating mechanism the motion is usually transmitted to the electrode by a clutch, which releases and lets the electrodes slip together.

In the carbon arc lamp, the mechanism is "floating"; that is, the upper carbon, held by the opposing forces of shunt and series magnets, moves with every variation of the arc resistance, and so maintains very closely constant voltage on the arc. In the long burning luminous arc, as the magnetite lamp, the light comes from the arc flame, and thus constant length of arc flame is required for constant light production. The floating mechanism, which constantly varies the arc length with the variation of the arc resistance, has therefore been superseded by a mechanism which sets the arc at fixed length, and leaves it there until with the consumption of the electrodes the arc has sufficiently lengthened to cause the shunt coil to

operate and to reset the arc length. Thus in some respects, these lamps are shunt lamps.

During the early days of the open carbon arc lamp, 9.6, 6.6 and 4 amperes were the currents used in direct current arc circuits, with about 40 volts per lamp. The 4 ampere arc very soon disappeared, as giving practically no light.

In the enclosed arc lamp, the carbons are surrounded by a nearly air tight globe, which restricts the admission of air and thus the combustion of the carbon, and so increases the life of the carbons from 8 or 10 hours to 70 to 120 hours. In these lamps, lower currents and higher arc voltages, that is, longer arcs, are used: in direct current circuits, 6.6 amperes and 5 amperes, with 70 to 75 volts per lamp; in alternating current circuits, 7.5 and 6.6 amperes are used with the same arc voltage.

In the direct current magnetite arc lamp, 4 amperes and 75 to 80 volts per lamp are used; in the alternating current titanium carbide arc lamp, only 2.5 amperes and 80 to 85 volts per lamp are used.

APPENDIX I

LIGHT AND ILLUMINATION

Paper read before the Illuminating Engineering Society, December 14, 1906.

REVISED TO DATE.

I.

COMPARED with other branches of engineering, as the transformation of electrical power into mechanical power in the electric motor, or the transformation of chemical into mechanical energy in the steam engine, we are at the disadvantage when dealing with light and illumination, that we have not to do any more strictly with a problem of physics, but that we are on the borderland between applied physics, that is engineering, and physiology. Light is not a physical quantity, but it is the physiological effect exerted upon the human eye by certain radiations.

There are different forms of energy, all convertible into each other, as magnetic energy, electric energy, heat energy, mechanical momentum, radiating energy, etc. The latter, radiating energy, is a vibratory motion of a hypothetical medium, the ether, which vibration is transmitted or propagated at a velocity of about 188,000 miles per second; and it is a transverse vibration, differing from the vibratory energy of sound in this respect, that the sound waves are longitudinal, that is, the vibration is in the direction of the beam, while the vibration of radiation is transverse.

Radiating energy can be derived from other forms of energy, for instance, from heat energy by raising a body to a

high temperature. Then the heat energy is converted into radiation and issues from the heated body, as for instance an incandescent lamp filament, as a mass of radiations of different wave lengths, that is, different frequencies. All kinds of frequencies appear: from very low frequencies, that is only a few millions of millions of cycles per second, up to many times higher frequencies. We can get, if we desire, still very much lower frequencies, as electromagnetic waves, such as the radiation sent out by an oscillating current or an alternating current; but the radiations which we get from heated bodies are all of extremely high frequency, compared with the customary frequencies of electric currents. At the same time they cover a very wide range of frequencies, many octaves, and from all this mass of radiations, from all the frequencies of radiating energy, somewhat less than one octave can be perceived by the human eye as light.

Light, therefore is the physiological effect exerted upon the human eye by a certain narrow range of frequencies of radiation. Frequencies lower than those visible to the eye, and frequencies higher than those visible to the eye, are again invisible.

We frequently speak of those frequencies which are lower than the visible ones, as radiating heat, and of those frequencies higher than the visible ones as chemical rays. This, however, is misleading, and there is no distinction in character between radiations of different frequency. There are no heat rays differing from light rays or chemical rays. Any form of energy when destroyed gives rise to an exactly equivalent amount of some other form of energy. If therefore we destroy radiating energy by intercepting the beam of radiation by interposing an opaque body in its path, then the energy of radiation is converted into some other form of

energy, usually into heat. That means that any radiation when absorbed produces heat and the amount of heat produced merely represents the amount of energy which was contained in the radiation. If the radiation contains a very large amount of energy, the heat evolved by intercepting it may be sufficient to be felt by putting your hand in the beam. If the amount of energy is less, it may not be possible to feel it, though with a sensitive instrument, as a bolometer, we may still be able to measure the heat. All radiations therefore are convertible into heat: the visible light waves as well as the invisible ultraviolet rays, and the—usually more powerful—long ultrared waves; but none of the radiations can be called heat, no more than the mechanical momentum of a flywheel is heat, because when destroyed, it produces heat.

If we consider the infinite range of radiation issuing from heated bodies, we find that those rays which are of lower frequency than the visible rays will be felt as heat, because they contain a very large amount of energy. The rays which are visible represent very little energy—and therefore they do not give as much heat. For instance, in the case of a hot steam boiler, although we get no light, we can feel the radiation from it by the heat which it produces when intercepted by our hand held near it. We do not feel the radiation as heat which issues from the green light of the mercury lamp, merely because the energy of radiation in the latter is less than the amount of energy in the radiation from a hot steam boiler; but while it is less in the former case, it happens to be of that frequency which affects the eye and is visible.

As a consequence, when we speak of cold light, this does not mean that it is different from hot light—from the light, for instance, given by a hot coal fire, where we feel the radiation as heat; it merely means that what is usually called cold

light (as the light of the firefly is supposed to be) is radiation containing to a very large extent rays of the visible frequencies and not much energy outside of the visible range; i. e., containing very little total energy, so that the energy when destroyed, that is, converted into heat, cannot be felt easily, but requires more delicate methods of determination; while a very inefficient light, as a coal fire for instance, which gives most of its energy as invisible radiation of low frequency, very little as visible radiation, can be felt by the heat produced by the interception of the rays, mainly the energetic low frequency rays. As stated, then, there is no essential difference between so-called heat waves and light waves, but any radiation can be converted into other forms of energy, the so-called chemical rays of ultraviolet light, the X-ray, as well as the ultrared and the visible rays, and when converted into heat can be noticed as such. Now it just happens that most of our means of producing radiating energy give high intensities of radiation only for very low frequencies, invisible ultrared rays, but we are not able to produce anywhere near the same intensities of radiation for higher frequencies.

So also, when we speak of ultraviolet, or short, high frequency waves, as chemical waves, that does not mean that they have a distinctive character in producing chemical action—any form of energy, naturally, can be converted if we know how, into chemical energy, the long ultrared waves just as well as the short ultraviolet waves. It just happens that those chemical compounds which are easily split up by radiating energy, are silver salts or salts of gold and platinum; they are especially affected by the ultraviolet and violet rays. We observe, then, the chemical action of these rays, but do not observe so well the chemical action of other rays. There may, however, be some feature in the constitution of matter, which

accounts for the high chemical action of the ultraviolet and violet rays.

It is obvious that if we intercept and destroy radiations and so convert their energy into other forms of energy, if the energy is only great enough, we get a high temperature, and thus a high chemical action, merely by the effect of temperature. But we may also get a chemical effect by what probably is some kind of a resonance phenomenon. The particles of a body, atoms or molecules, must have some rate of vibration of their own. If, then, a ray of radiation impinges upon them which is of a frequency of the same magnitude as the inherent rate of vibration of the atom, by resonance this vibration of the atom must rapidly increase in intensity until the atom breaks away from the others, or the molecule breaks up, that is, the chemical combination is split up.

The inherent frequency of oscillation of the atom seems to be of about of the same magnitude as the visible radiation, or rather of a little higher frequency; that is, if the atoms are left to vibrate freely as under the influence of an electric current in the arc, then we get radiations of the frequency inherent to the atom. The general tendency then is toward the violet or short wave end of the spectrum. If we assume that the mass of the silver atom is such as to give a rate of vibration in the range of the violet and ultraviolet, it is easy to understand that radiation of this frequency splits up the silver salt by increasing the vibration of the atom by resonance, and that shorter or longer waves have no effect, or much less effect. So it may be a mere incident that those chemical compounds on which we observe the chemical action of radiation just happen to be sensitive to the violet end of the spectrum. It is indeed a fact that other chemical changes brought about by radiating energy, as the formation of ozone

from oxygen, that is, the splitting up of the oxygen molecule and reforming of the ozone molecule from the atoms, do not take place in the violet or ultraviolet, but requires frequencies very much higher, about the highest frequencies which the mercury arc at low temperature gives. Possibly, since the oxygen atom is so much lighter than the silver atom, its frequency of vibration is much higher, which means that resonance effects and destruction of the molecules take place only with a much shorter wave length of radiation, or much higher frequency.

Vice versa, it seems that these frequencies which are chemically active in organic life, which give the energy absorbed from radiation by plants, and so the chemical activity utilized in building up the growth of vegetation, are not at the violet, but at the red end of the spectrum. It appears that the red and ultrared rays produce growth of plants and the chemical activity which we call life, while the violet and ultraviolet rays kill. Even this we can well understand if we consider the chemical activity as a resonance phenomenon, because the metabolism of protoplasm which we call life, is based on the existence of unstable structures of carbon compounds. We have here not atoms combining with each other, but groups, chain and ring formations, which are of larger mass and therefore have a lower rate of vibration and so should be expected to respond to lower frequencies or to red light, as indeed seems to be the case. The violet and ultraviolet light does not split up the organic matter into groups, which recombine to form complex bodies, and so represent the changes called life; but due to its higher frequency, resonance occurs with the atoms, that is, the organic compound splits up into atoms and so disintegrates, which means death.

So it can be understood that the chemical activity of different radiations may be different; the chemical activity of long rays gives life to the vegetation and the short waves,

The popular distinction between heat waves, chemical waves and light waves, therefore is not a physical distinction, destruction down to the atom.

death; one splits up into carbon groups and the other carries but all are radiating energy of the same character, differing merely in wave length, and the visible range is somewhat less than one octave, rather at the upper end, at the higher frequencies, which are difficult to produce. This makes the problem of investigating and dealing with light difficult for the engineer, because it is not any more a physical quantity which can be measured accurately, as in the case of power or velocity, but it is a physiological effect. We can, indeed, measure very accurately the total energy of radiation from a heated body, but the total energy of radiation is not light: only a very small part of it is visible. We can go further and split up the total radiation issuing from a hot body, as the incandescent lamp filament, into its different wave lengths and different frequencies; as for instance, we can resolve the total radiation into the spectrum by using a prism to separate the different frequencies, and then collect the total of the radiation within the visible range, by a lens or other means, and measure all the energy of the visible radiation. Or, still simpler, although approximate, we may interpose in the beam of radiation some medium which absorbs the invisible long rays and invisible short rays, and which transmits, all or rather most, of the visible rays, as for instance glass and water. In this manner one could easily measure the energy of the visible radiation, and compare the energy of the visible radiation with the total energy producing this radiation. That would give a physical measure of the

efficiency of producing visible radiation but it would not be a measure of the efficiency of producing light, since unfortunately the different wave lengths of visible radiation are very different in their physiological effect. The same amount of energy as visible radiation, giving the effect of green light, represents an entirely different amount of light, a many times greater physiological effect than the same amount of energy as red rays, that is, rays of the wave lengths which give the impression of red light.

That means, the physiological effect or light-equivalent of mechanical energy within the visible range—is a function of the wave length and varies with the wave length, that is, with the color. That really is obvious, if you think of it: if you follow the range of frequency from a low frequency to a high frequency, you see that energy radiating at low frequency represents no light whatever, has no physiological equivalent, is invisible. When you come into the visible range it has a physiological effect. Therefore, when you pass from the invisible into the visible range, the physiological equivalent must pass from zero into a finite value and must necessarily pass continuously, that is, at the extreme end of the visible range; the light equivalent of energy must be extremely low, and the further you go into the visible range, the greater it is, reaching the maximum in the middle of the visible range, in the green and yellow, and decreasing again down to zero at the violet end of the visible rays; beyond that, at still higher frequencies, the physiological equivalent of energy is zero again; or, vice versa, if we consider the mechanical equivalent of light, it is a minimum in the middle of the visible range, where one candle power of light represents the lowest amount of energy, and increases toward the ends of the spectrum of the visible range, to infinity at the ends of the visible range.

Now, that means, in plain language, that the efficiency of light production is a function of the wave length, that is of the color, and that it is at its maximum in the middle of the spectrum, where the same amount of power, measured in watts, gives the largest amount of light measured in candle power. So the efficiency of light production is a function of the wave length.

Unfortunately, the physiological equivalent of power, or the physiological effect of light varies not only with the wave length, but also with the absolute intensity. Suppose we undertake to compare red, yellow and green lights, or any lights of different colors. First we meet great difficulties in comparing them. We want one candle power in light, as red, yellow, or green. You cannot compare different colors of light directly, since there is no physical measure of light. Lights are compared with a standard lamp, which has a certain color, yellowish white. A light of the same color we can compare exactly; if the color is not much different, we still get an approximate comparison; but with widely different colors, we obviously can not get even an approximate comparison, can not say when the two sides of the photometer screen, one illuminated by green light, the other by red light, are equal in intensity. There is thus no direct comparison of differently colored lights. You have then to go one step farther and consider that light is used for illumination, is used to see by, and this gives you a fair comparison: you observe at what distance from the two lights, red and green, you can read with the same convenience, read the same kind of print, or to measure more exactly, get the maximum distance at which you can just read a certain size of print, by either light. At that distance the two illuminations are the same, and the two lights so have an intensity inversely proportional to the square of

these distances. In this manner lights of different color can be compared.

Necessarily, the comparison has not the accuracy of photometrical comparison. This cannot be expected, since you do not compare physical quantities, but only physiological effects on the eye, and different observers may have different personal contents. The eye of the one may be more sensitive to green, and the eye of the other may be more sensitive to red, and therefore the comparison may be different. However, these individual differences are not great, and different observers, even with widely different colors of light, do not give results differing much from each other, so that a comparison of intensities of differently colored lights, and thereby a measurement of the intensity of differently colored lights in candle power is feasible, by some such method, that is, of observing the illumination produced by the different lights.

You find, however, if you have a green light and a red light, which at a certain distance appear equally brilliant to the eye, then when you get nearer to the two lights the orange red light appears much brighter than the green, and when you go further away the green light appears brighter, and at still greater distances you still see the green light fairly brightly, while the red light is almost invisible. That is, the relative physiological effect of different wave lengths varies, not only with the wave lengths, but also with the absolute intensity of illumination, and while throughout the whole range the sensitivity of the eye for green light is much greater than for red light, the difference is far greater for low than for high illumination, that is, the ratio of sensitivity for green compared with that for red is greater for faint illumination than for intense illumination. If you desire to express lights of different colors in candle power it therefore seems necessary

also to give the distance, or the intensity of illumination at which you have observed; in other words, the light from the middle and the short wave end of the spectrum gives a better and more efficient illumination where the total intensity of illumination is low, while the long wave or low frequency of the red and orange and yellow light gives a much more brilliant effect at high intensity than the same volume of light of shorter wave length.

This is of importance for the illuminating engineer, because where you desire to get high intensity effects, as in decorative lighting or in advertising, better results are given by the low frequency end of the spectrum, by orange and yellow light, whereas when you are satisfied with low intensity of illumination, as in street lighting, you get better results from the short wave end or the middle of the spectrum, from the greenish-yellow of the Welsbach gas light and the bluish-green of the mercury lamp, and not from the orange-yellow of the old incandescent lamp. Therefore the white light of the carbon arc gives better results in street lighting than the yellow of the incandescent lamp, even at equal intensity of illumination. These features have been of less importance until a few years ago, since the available sources of light were all approximately of the same color, varying from the orange-yellow, to yellow and yellow-white, to white; from the gas lamp, kerosene lamp and tallow candle of orange-yellow color, to the yellow incandescent lamp and the yellowish-white arc, yellowish-white sunlight, to the white diffused daylight. This was a fairly limited range. It is only in the last few years that illuminants of high efficiency have been brought out, which give marked and decided color differences, and are available in units of suitable size and of high efficiency, as the greenish-yellow of the Welsbach gas lamp, the bluish-green of the mer-

cury lamp, and the orange-yellow of the flaming arc, and hence these questions are increasing in importance.

II.

This brings us to the consideration of the methods of producing light. Until a few years ago, until the development of the Welsbach gas mantle, practically all methods of producing light were based on incandescence. That is, by impressing energy on a solid body, either the chemical energy of combustion, or electric energy in the incandescent or carbon arc lamp, the temperature is raised to such a high degree that amongst the total radiation issuing from the heated body a certain very small percentage appears within the fraction of an octave of visible light. With increasing temperature of the radiating body, the average wave length of radiation decreases, that is, the average frequency of radiation increases and so approaches nearer to the visible range, although still at the very highest temperature which can be produced the average wave length of radiation is very far below the visible. This means that the higher a temperature is reached by an incandescent body, the higher is the average frequency of radiation, and therefore the larger a percentage of the total energy of radiation is within the visible range, as light. The problem of efficient light production by incandescence therefore is the problem of reaching as high a temperature as possible in the luminous body. In the gas flame and the kerosene lamp, this temperature is the temperature of combustion, rather limited. In the incandescent lamp it is limited also. In the latter case the temperature which can be reached is limited by self-destruction of the incandescent body.

The highest temperature probable which can be reached is the boiling point of carbon; it is reached in the crater of the

carbon arc lamp, and therefore the carbon arc gives the most efficient incandescent light. It is incandescent light, because it comes from the incandescent crater, and the arc flame or the vapor conductor does not appreciably contribute to the amount of light issuing from the arc lamp. Very much lower, necessarily, is the temperature of the incandescent lamp, of the carbon filament.

The problem is to find materials which can stand very high temperatures, to increase the temperature of the gas flame as well as of the incandescent filament. We have increased the temperature of the gas flame by using a gas of higher chemical energy, as acetylene. The acetylene flame is white; the ordinary gas flame is yellow. We have increased the temperature of the carbon filament by replacing the carbon with some more refractory material, such as tantalum, osmium, tungsten, etc., and thus getting a higher efficiency. We can increase the temperature of the gas flame by increasing the rapidity of combustion. We can increase the temperature of the carbon filament in the incandescent lamp by increasing the energy input, but if we increase the temperature of the carbon filament, it is more rapidly destroyed. If we increase the temperature of the gas flame by more rapid combustion—to a certain extent we have done it already, by having the gas issuing not from a round hole, but from a flat slit, so as to give a larger surface to the flame; if we go still further and mix the gas with air, we get a still higher temperature, a more rapid combustion, but we lose the incandescent body, because the incandescence of the gas flame is the light given by carbon or heavy hydrocarbon particles, floating in the gases of combustion. We could increase the efficiency of the gas flame by mixing the gas with air, as in the Bunsen flame, but we have then to insert a luminous body of some other material, as no carbon is pro-

duced by the gas in its dissociation. We can do it by a skeleton of platinum wire. In no case, however, can we reach very high efficiencies by incandescence, due to the temperature limit.

We could, however, increase the efficiency of light production if we could find an incandescent body which would not radiate in the same manner as the carbon filament or the so-called black body, but which would give an abnormally low radiation in the low frequency range, or an abnormally high radiation in the high frequency or visible range. Such a body may be said to give selective radiation, because the distribution of energy in the spectrum amongst the different frequencies of radiation is not the same as it would be with an ordinary black body of the same temperature. If we found a body which would give an abnormally low radiation in the visible range, or abnormally high radiation in the invisible range, this body would be an abnormally inefficient light producer. Vice versa, if we found a body giving abnormally high radiation of short wave lengths, in the visible range, or abnormally low radiation of long waves, of low frequency, this would give an abnormally efficient incandescent body. Such bodies exist and the enormous progress in gas lighting made by the introduction of the Wellsbach mantle is based on selective radiation, that is, the oxides do not radiate the same range and intensity of waves as a black body, the incandescent carbon, but give an abnormally large amount of visible rays compared with invisible rays; that is to say,—they give a larger percentage of high frequency light rays compared with the low frequency invisible rays. Possibly and even probably some of these highly efficient filaments like the tungsten filament, also owe some of their high efficiency of one watt per candle power to selective radiation.

When discussing selective radiation, we have first to come to an agreement on what we understand by selective

radiation. The question whether an illuminant owes its high efficiency to selective radiation, depends largely on the definition of the term "selective radiation". We have here a similar case to that of the much discussed problem of the "counter electromotive force of the electric arc". Whether the electric arc has a counter e. m. f. or not, entirely depends on the definition of counter e. m. f. In the same way, the decision on the question of selective radiation depends upon what you define as selective radiation. If you define as selective radiation any radiation in which the intensity of radiation is distributed through the total spectrum differently from that of the theoretical black body, then the Welsbach mantle has selective radiation. If, however, you define selective radiation as the radiation of a body which gives spectrum lines, or bands, or absorption lines and bands, that is, sharply defined narrow ranges in the spectrum, of abnormally high or abnormally low intensity, then the Welsbach mantle has no selective radiation. So all discussions and statements on selective radiation have rather little meaning, if the writer does not give his definition of selective radiation. In the following, I define as selective, any radiation which differs in the distribution of its intensity from the radiation of the theoretical black body.

In an incandescent lamp filament we do not get a definite pitch, or definite frequency of vibration, but we get an infinite number of different waves. The reason is perhaps, that in a solid or liquid body the vibrating particles are so close together as to interfere with each other. If you could set a body in vibration, in which the vibrating particles, atoms or molecules, are so far apart as not to interfere with each other, as in a gas at low pressure, then they would execute their own periods of vibration, and then the light from such a body would not be a radiation of all wave lengths, but we would get radiations of

only a few definite wave lengths, or a line spectrum. So incandescent or luminous sodium vapor gives only one kind of light, a yellow spectrum line, and in addition thereto a number of utrared and ultraviolet rays.

Since the spectrum light is based on the non-interference of the vibrating particles, it is easy to understand, that when you bring the atoms or molecules closely together—as at atmospheric pressure—interference may begin, and the lines of the spectrum become more confused and blurr into bands. Therefore, we see in the mercury arc spectrum, which is at low vapor pressure, a small number of definite, sharply definite lines. In the calcium spectrum of the flame carbon arc, we get a large number of lines blurring into each other to an almost continuous spectrum; so also in the white spectrum of the magnetite-titanium arc.

When we set a gas or vapor in vibration, it vibrates at its own frequency, independent of the temperature, and it is merely a question of the character of the material whether a very large percentage of the total energy of radiation happens to be within the visible range or outside the visible range. Temperature does not come in as factor, because the frequency of radiation is no longer a function of the temperature, but independent of the temperature. Sodium vapor gives the same frequency of radiation, the same yellow line when the sodium vapor is at low temperature or at high temperature. Some spectrum lines may increase in intensity with an increase of temperature faster than others, and the color of light may change with the temperature, change from yellow to white or blue, or from green toward white, and red, as the mercury light does with increasing temperature, but that is merely a characteristic feature of that particular body, and not a general character of the temperature effect; the possibility there-

fore exists of finding materials which, when energized, as vapor or gas, give a spectrum with a large amount of energy in the visible range, thus giving an efficiency of light production far in excess of that available by incandescence.

So far the only materials which give these characteristics are mercury, calcium and titanium. These three metals give spectra which contain such a large percentage of the total radiation in the visible range, that the amount of light measured physiologically in candle power is far in excess of that which possibly could be produced by incandescence, even with the assistance of selective radiation. Their industrial applications are represented by the mercury arc, the yellow flame carbon, and the white magnetite and titanium arc, and these are of such very high efficiency as to be of higher magnitude than any incandescent light.

Even if we consider only these three illuminants, we have quite a color scale. From the orange-yellow of the flame carbon, which is about the longest wave length we could use, to yellow and yellow-white, in the acetylene flame and the tungsten filament. Then we have the greenish-yellow of the Wellsbach mantle, by selective radiation. We have the bluish-green in the mercury arc, and the yellowish-white of the carbon arc, as well as the clear white of the titanium arc. Each of these can be modified. We can modify the titanium arc, giving all colors from yellow-white to bluish-white, by the addition of other materials which give either yellowish or bluish spectra. We can modify the yellow calcium arc, from the orange-yellow of calcium fluoride down to the yellowish-white of calcium borates. You can modify each color over a certain range, and you can get pretty nearly any color, with the exception perhaps of a clear blue and violet: no means have been found to get approximately the same efficiency in those

colors of very short wave length, as in the other colors of lights.

This feature makes the effect of color which I discussed before, the variation of the physiological effect with the brilliancy of illumination, of more importance now than years ago, when the only method of producing colored light was by the absorption of all other colors.

III.

After all, however, it is not light, that is wanted, but illumination; it is not the amount of visible rays issuing from the source of light, the incandescent lamp or gas flame, which is of importance, but the amount of light which reaches the objects we desire to see, that is, the illumination produced by light. In this respect, I believe a mistake has been made by the gas industry as well as the electric lighting industry, for many years, by devoting all energy to the production of light, the development of the lamp, while they have almost entirely left out of consideration, that the production of an efficient light is not the only important problem, but that of the same importance is the arrangement of the light so as to get efficient illumination, that is, get the greatest benefit from the light produced; and this feature has been usually left to the tender mercies of the architect or the decorator, who placed the lights wherever he thought they would look artistic, regardless of the requirements of effective illumination. If you look around, you find cases everywhere of artificial illumination where the lights have been arranged, so that you get a very poor illumination from a large amount of light. To overcome these defects, it is necessary to study the problems involved between the production of light and the physiological effect produced by the light upon the eye, and it requires a careful study, just as

any other engineering problem. It is of importance to consider not only the amount of light issuing from the source, but the amount of light which reaches the object to be seen by the illumination.

The demands of illumination are mainly of two classes, *general illumination*, and *local*, or *concentrated illumination*. Many cases require general illumination, as a meeting room, where it is desired to see equally well everywhere; that is, to get the same intensity of illumination throughout the whole illuminated area. So also a draughting room, a school room, the hall of a house and the streets of a city require general illumination, a uniform fairly high intensity in a draughting room or school room, a relatively low but as nearly as possible a uniform intensity in the streets of a city. It is true, street lighting is usually very far from uniform, but that merely means that the problem of proper street lighting is usually not solved in the most efficient and satisfactory manner. In other cases concentrated lighting is required, as in domestic lighting, in the dining room, the living room, etc., where light is desired on the table where we work, eat, read, etc. In such cases, the general illumination of the room is of lesser importance; it is not needed to any extent, or is frequently undesirable, because a room with a very low intensity of general illumination frequently is considered more homelike, especially by the feminine part of the human race. In still other places general illumination may be directly objectionable, as in a sick room. Most cases, however, require a general illumination of moderate intensity, and a far more intense local illumination, as over the desks in an office, or the reading tables in a library. In such cases merely a general illumination would be sufficient, if very intense, but this is uneconomical and to some extent objectionable on account of the blinding glare, which is disa-

greeable; and so a combined general and local illumination is more efficient and more satisfactory.

In producing illumination either *direct lighting* or *indirect lighting* may be used. That is, the rays issuing from the source of light may either pass directly to the illuminated objects, or they may pass to a reflecting surface, and be reflected from this surface to the object, or may pass through a refracting body, as the frosted incandescent lamp globe, or opal globe of the arc lamp, and so reach the illuminated object. In general, it is obvious that any method of indirect lighting by refraction or reflection wastes a considerable amount of light. That means, the total amount of light which reaches the illuminated object must necessarily be less with indirect lighting, as compared with direct lighting, with the same amount of light.

Indirect lighting can be done by reflection or refraction by some attachment to the lamp, as a reflector or a holophane or frosted globe, or by reflecting the light from the ceilings and walls of the room, on the objects to be illuminated. In the latter case, it is obvious that white walls give the highest efficiency of reflected light. It is easy to observe that the same source of light in a room with white walls gives several times the intensity of illumination which it gives in a room with black or non-reflecting walls. That means that the total amount of illumination is increased several-fold by reflection from white walls. So in a draughting room, or school room, by using as light walls as possible, we get the best efficiency of illumination.

It is not always feasible to have light walls, especially when you come to machine shops or foundries, and other places where the walls do not remain white, but change to some darker color. The question is, what color do these walls

assume? The color of almost everything which is changed by age is due to either iron or carbon. In most cases of discoloration by age you see the reddish-brown color of iron and the brownish-yellow color of carbon. This is the color most subjects gradually assume. This color of age is in the long wave or low frequency end of the spectrum. To get the benefit of reflected light from walls which cannot be kept perfectly white, a source of light rich in the long low frequency waves, or of a yellowish tinge is therefore more efficient by giving more reflection from the walls than a source of light rich in short or high frequency waves, that is, bluish-white. This effect is very marked when you compare the mercury lamp with the flame carbon lamp. The illumination given by the mercury lamp in a draughting room is very satisfactory. The same illumination in a foundry or machine shop is far less satisfactory, and you notice there a marked absence of reflected light, that is, the walls and ceilings all gradually assume a color which is rich in red and yellow, and so reflect very little light of the violet end of the spectrum. Even the black-begrimed walls of a blacksmith shop reflect a considerable amount of light with an orange-yellow source of light; practically none with the bluish-green of the mercury lamp.

Thus the shade of color of the illuminant may be very essential in getting efficient illumination. In the interior of a city, the walls usually have a reddish-yellow color. In that case white or yellowish lights are superior. When outside of a city, the greenish-yellow of the Welsbach lamp, the bluish-green of the mercury arc, gives a much larger amount of reflected light from vegetation than the yellow of the incandescent light and so a better illumination. Vegetation absorbs the long waves, the low frequency radiation; so with a yellow source of light there is practically no reflection from

living vegetation, but only reflection from dead vegetation, and in the light of the incandescent lamp or the flame arc all vegetation appears very poorly, the dead parts are very prominent, while the reverse is the case where the light is deficient in the red and yellow, and rich in the green and blue, as with the mercury arc: the green shows prominently, while the dead leaves, etc., are not visible.

IV.

It is, however, not the amount of light which reaches the illuminated objects which is of importance, that is, not the physical intensity of illumination, but the amount of light which from these illuminated objects reaches the human eye. With the same intensity of illumination, the same amount of light reaching the illuminated object, of the same color, the amount of light entering the eye may vary widely with the opening or contraction of the pupil of the eye. The eye is automatically adjusting for intensity of light. This is the reason we see well at sun light and at the light of the full moon, although the former is many thousand times greater in intensity than the latter. The eye can accommodate itself to intensities varying over an enormous range. It does this partly by the fatigue of the nerves of vision, partly by the contraction or opening of the pupil. This is undoubtedly a protective device developed in the human race. It means that if we have in the field of vision a source of light of high intrinsic brilliancy, the eye protects itself by contraction of the pupil and so it receives very much less light in the field of vision where we want to see objects, than if the source of light were taken out of the field of vision. By eliminating the source of light from the field of vision and eliminating the contraction of the pupils resulting from the high intrinsic brilliancy of the illuminating

body, we get actually a much larger amount of light into the eye with the same amount of light striking the illuminated object; that is, we get a higher physiological efficiency. Even with a much smaller amount of light reaching the illuminated objects, we still get more light in the eye. That means if we reduce the intrinsic brilliancy of the illuminant by indirect lighting, by diffusing the light, we may lose a considerable amount of light, actually get a considerably reduced quantity of light on the object which we desire to see, but still we get a larger amount of light from these objects into the eye, because the eye is open further and admits more light and is less fatigued.

It follows from this that in efficient illumination, it is of foremost importance to arrange the illuminants so as not to have excessive intrinsic brilliancies in the field of vision when looking at the objects we desire to see. That means that the proper field for the illuminant is outside of the field of vision, or where you cannot get it out of the field of vision, that its intrinsic brilliancy should be reduced by diffusion: thereby we actually get a much higher physiological effect. This is the reason for indirect lighting. We may have a very large amount of light thrown on any object in a room, but if the eye is fatigued by seeing the source of light in the field of vision, we get very little light in the eye, while with a properly arranged indirect lighting, with a much lower amount of light reaching the object, we get a higher physiological effect, that is a better and more efficient illumination.

It appears, however, that this automatic protective faculty of the eye was developed through the ages as a protection not against light, but against energy; apparently the eye is protecting against the energy of radiation, not the physiological intensity, and since the energy of radiation is mainly in the

ultrared, in the long waves, the frequency which causes the protective reaction is the frequency of the long wave end of the spectrum, the red and yellow waves; they make the pupil contract. This action is much less for the green and blue rays. That is the reason the eye does not react on the mercury lamp to any great extent. It means a green light, like the mercury or Welsbach lamp, can be in the field of vision to a much greater extent without causing the contraction of the pupil and so reducing the physiological effect. This is of importance in places where the light cannot well be taken out of the field of vision, as in street illumination. In this case, all the sources of light must be arranged along the street and so must be in the field of vision. By cutting off the red end of the spectrum you eliminate the contraction of the pupil, and get the full benefit of the light between the illuminants, while with a yellow source of light, as with the incandescent or arc lamp of old, you do not get the benefit, that is, the physiological effect of the illumination by a green illuminant in such cases is superior to that by a yellow illuminant, the illumination appears brighter and more uniform.

A light devoid of red and yellow rays is at the same time the safest and most harmless, and also the most harmful. It is the safest and most harmless, and gives the most uniform illumination, if its intrinsic brilliancy is sufficiently low to be below the danger limit of energy of radiation, but it is harmful if above that, because the eye does not protect itself against it, probably because these lights have not existed throughout all the ages when this protective action of the eye was developed, and sunlight and fire were the only sources of light, both rich in red rays. This accounts for the rather contradictory effect observed, that green or blue light, as the Welsbach mantle or mercury lamp, is a very good light to work by,

superior to the yellow kerosene lamp, and at the same time there is some suspicion that it is harmful to the eye. It may well be that where it is of very high intensity, the automatic protection of the eye is not sufficient to protect with such light. Where you use such sources of light you can get the benefit of the absence of the contraction of the pupil, but it devolves upon you to arrange the illumination so as not to get the harmful effects against which the automatic protection of the eye fails. That means all these lights are superior for illumination if they have low intrinsic brilliancies, but somewhat questionable if they have extremely high intensities.

V.

It is, however, not even the amount of the light which enters the eye which is of importance in illumination, but the difference in the amount of light. If in the illuminated area the light were of uniform intensity, and everything of the same color, we would see nothing but a glare of light. The seeing takes place by a difference in color, and difference in intensity. Difference in intensity includes shadows. Shadows are thus an essential feature in seeing things.

Considering, then, the seeing by shadows and seeing by color differences, you observe that by this feature we can divide illumination into *directed* and *diffused illumination*. In diffused illumination light comes in all directions with approximate uniformity, and shadows do not exist; in directed illumination, shadows exist. In some cases shadows are objectionable, and in other cases shadows are necessary for clear distinction, and diffused illumination in such cases would not be satisfactory.

As regard to seeing differences in color, it is obvious that where definite color distinctions are required, you can intensify

the sharpness of vision by selecting the color of your light best suited to bringing out the colors desired. Where the color conditions you want to distinguish are those due to age, iron and carbon, then the light which is deficient in red and yellow, which therefore shows the colors given by iron and carbon, as black, gives a much sharper distinction, and the mercury lamp shows blemishes and dirt much more pronounced than the white light. Again, the sources of light which are very rich in red and yellow rays show these colors due to iron and carbon very much less, and therefore show blemishes or a slight amount of dirt much less, soften them; and where the color distinctions are those due to these two most prominent elements, in the yellow light their appearance will be greatly softened, and under the green light they will be made harsh and sharp. If you desire to soften effects, as in a ballroom, it would be fiendish to use mercury lamps, but where you want to search out a spot that is soiled, it would be very wrong to use a dull, yellow incandescent lamp or a gas flame, but rather to use the green Welsbach light, or better still the bluish-green mercury arc, which gives in such case sharp distinction, where white light shows little, and yellow light nothing. Where you desire to see all colors in about the same relation as by daylight, you obviously desire a white light.

It is therefore important for the illuminating engineer to select the shade or color of the light and study the requirements of each case which comes into his charge. It would be just as wrong in one case to use an incandescent lamp, where the mercury lamp would be better, as to do the reverse.

We have to distinguish then between general illumination and local illumination, between direct illumination and indirect illumination, and between directed illumination and diffused illumination. These three different classes or distinc-

tion to a certain extent overlap. It would be very wrong, however, to mistake them, and a very serious mistake in the design of a system of illumination can very easily be made; for instance, by mistaking general illumination and diffused illumination for each other. The problem may be to get uniform intensity all over. You can get that by distributing a large number of small units all around the cornices and reflect the light from white walls and ceilings and get a very diffused illumination, or you could get general illumination, where the intensity of illumination all over is the same, in a moderate sized room, from one source of light by using one of these sources of light as an incandescent lamp with a holophane reflector, which gives the proper distribution, or you could get light from any other source by controlling the distribution curve of the light so as to get uniform distribution. The former arrangement gives diffused light, the latter directed light. You may get the same intensity of illumination all over the room, in both cases; but in the former case no shadows, in the latter case absolutely black shadows. Probably in the former case for domestic use the lighting will be unsatisfactory and trying to the eyes, because you do not see well, you do not have any shadows, objects around you are not so distinct, because you lose the distinguishing feature of the shadow. In the latter case with the directed lighting from one source, the lighting will be unsatisfactory because you get very dark shadows, and you do not see anything in the shadows, and the eyes will be made tired by trying to see in the very dark shadows.

You have to consider how much directed light and how much diffused light you require. In some cases you may desire only diffused light. In the general lighting of a draughting room you do not want any directed light, since

you must have no shadows, because if the ruler casts a shadow, it is trying to the eyes to distinguish between the edge of the ruler and the edge of the shadow, and mistakes may be made. In this case, you see only by differences in color and in the intensity, and not by shadows. You therefore get satisfactory illumination from many small units, or by indirect lighting, reflected light from white walls and ceilings, but you get unsatisfactory illumination from a few units even when properly distributed so as to give uniform intensity all over, but giving little reflected light. In other cases, you may also require a general illumination equal in intensity all over, but you need directed illumination so as to see by the shadows. So for instance, a good draughting room illumination would not be suitable for a foundry. In a foundry, where all the objects assume more or less the same color, you require shadows to see by. Then you need a number of units of light to give directed illumination, but you must not go so far as to be unable to see in the shadows; you must have some diffusion, or overlapping of the different beams of light. So if you take a satisfactory foundry illumination and put it in the draughting room, even if the intensity were satisfactory, it would be entirely unsatisfactory, and so would be the reverse. It is therefore not merely the distribution of the intensity of the light, which is essential, but also the character, whether diffused or directed light, or how divided between diffused and directed light.

In the different lighting problems you therefore meet the question of concentrated and general illumination, of directed and diffused illumination. In domestic lighting, by reflected light from white walls and ceilings we can get a high intensity, and can increase the illumination several-fold over that given directly from the source of light, such as the incandescent

lamp or gas flame. Still the illumination would be unsatisfactory and tiring to the eyes. We all know that in the home a room with white walls is not as agreeable as one with darker walls. We say we have too much light. But we do not have too much light, because we do not have anywhere near the same amount of light as we get during the daytime out of doors. We have too large a percentage of diffused light. The intensity of diffused light is too great as compared with the directed light. We lose the shadows and that is tiring to the eyes. The problem of domestic lighting then is, to get sufficient directed and not too much diffused lighting so as to get the best vision, that is, to get sufficient shadows to see by, but the shadows must not be so dark as to make seeing objects in the shadows tiring to the eyes. During the daytime we get directed light from the windows, diffused light reflected from the walls. To get the proper proportion between directed and diffused light, fixes the shade of the walls, and in general we have to use walls of somewhat darker color. When you come to lighting in the evening, with a source of light like the incandescent lamp or gas lamp, sending out light in all directions, the diffused light compared with the concentrated or directed light is a higher percentage than in the daytime for the same color of walls, partly due to the color of the light, which is yellow, and is more reflected from the walls, largely, however, because with the daylight through the window the directed light is a much larger percentage of the total light than in the lamp, where only a small part is concentrated light. It is not comfortable to have this strong diffused light, and so we put shades on which absorb three-quarters of the light, but which give us a more comfortable illumination in the room. That means waste, however, and you pay for light which you do not use.

The proper illuminating engineering then is to secure the correct distribution curve of the source of light, so as to give the desired amount of concentrated lighting on the dining or reading table, and give only as much diffused lighting as is compatible with the amount of direct light used, to see in the shadows. The problem of domestic lighting, from the illuminating engineering point, is to determine the illumination over the entire area, and also the character of illumination, whether directed or diffused; how large an amount of light should be concentrated, and how large an amount should be directed; then the question of colors and shades also comes in as an important factor, as was discussed before. Practically nothing has yet been done in this direction systematically and intelligently, but all has been done by trial which at the best usually means producing more light than necessary, and throwing away the excess of diffused light by absorption.

APPENDIX II

LIGHTNING AND LIGHTNING PROTECTION

Paper read before the Annual Convention of the
National Electric Light Association, 1907.

Revised to date.

I. LIGHTNING PHENOMENA IN THE CLOUDS.

THE first man who attacked the problem of lightning and lightning protection, a century and half ago, was our great citizen, Benjamin Franklin. He gave us the lightning rod, which is now universally recognized as the most effective and only protective device for isolated points, as steeples, chimneys, etc. The next step in advance was made by Faraday: he showed that in the interior of a perfectly conducting body no electric disturbances can be produced by outside electric forces. This led to the most effective protection possible against lightning or electric disturbances, the use of a grounded metal cage, "Faraday's cage", enclosing the structure which is to be protected, whether a building against lightning, or a delicate instrument against electric fields.

In its simplest form, Faraday's cage, applied to a transmission line, is the ground wire above the line, and the protection afforded by it is the more complete, the more the overhead ground wires represent the condition of an enclosing cage of perfect conductivity. That is, a system of wires above and on the sides of a transmission line is superior to a single wire, a wire of high conductivity superior to a small iron wire. Here I specially desire to draw attention to the second requirement of the Faraday cage, high conductivity. Thus it is not sufficient merely to have any kind of overhead

grounded wire regardless how small, but high conductivity of the grounded conductor is essential in many cases of atmospheric disturbances.

In the last ten years, transmission voltages have crept higher and higher, transformers have been built of considerable size, of still higher voltages, so that exact data on the action of voltages up to 300,000 are now available, and approximate data for potentials above a million volts. It was found that air has a definite and fixed breakdown strength, that is, just as a beam breaks mechanically as soon as the stress in it exceeds a definite value, the breaking strength of the material, so air breaks down by a disruptive spark, as soon as the electric stress in the air, or the potential gradient, exceeds a certain value, which is about 100,000 volts per inch. The disruptive strength of air is, over a wide range, proportional to the pressure, that is, at a pressure of two atmospheres it is twice as high, or 200,000 volts per inch; at one-quarter atmosphere it is one-quarter, or 25,000 volts per inch.*

The striking distance in air between needle points has been investigated up to 300,000 volts, and found that for high voltages it is very nearly 10,000 volts per inch, that is, a discharge of 30" length between needle points requires 300,000 volts. If between two needle points the potential difference is gradually increased, already at relatively low voltages the disruptive strength of the air at the needle points is exceeded, the air at the points breaks down and becomes conducting, and luminous, as "brush discharge", so that the terminals are not **the** needle points any more, but the whole space, of approximately spherical shape, which is covered by the brush discharge. As result thereof, for high voltage, no appreciable

* Only at very low pressures, where the distance between air molecules become appreciable, this law ceases, and the disruptive strength increases again, and seems to become infinitely great in a perfect vacuum.

difference exists in the striking distance between needle points and between spheres, the centers of which approximately coincide with the needle points, as long as the diameter of the spheres is small compared with their distance apart. With increasing potential difference between needle points, the brush discharges spread out against each other, until only about 40% of the space between the needle points is free, and then a disruptive spark passes.

Naturally, as soon as determinations of spark voltages became available, attempts were made to estimate the voltage of a lightning flash. Considering, in a lightning flash, the discharge as that in an ununiform field, similar to that between needle points, and so requiring about 10,000 volts per inch. In this case, a lightning flash of two miles, or about 10,000 feet length, would require a potential difference of about 1200 million volts. The existence of such voltages in the clouds does not appear possible: a potential difference of 1000 million volts would produce a brush discharge of about one-half mile in length, before the final lightning flash occurs. In the brush discharge the air is electrically broken down, and becomes conducting. But it is also mechanically and chemically broken down, that is, the molecules are dissociated and recombine after the discharge, in all possible combinations. That is, we get ozone and nitric acid, and a lightning flash produced by a thousand million volts would thus be followed by a deluge of nitric acid. This fortunately is not the case.

An estimate of the voltage and the current in a lightning flash would not yet give the energy, if the duration of the discharge is not also known. We can, however, get an approximate estimate of the magnitude of the energy of the lightning flash indirectly, from photometric considerations, and eliminate the consideration of the duration of the flash by the

integrating feature of the human eye for impressions of very short duration: an impression on the human eye persists for some time, about .1 seconds, and any phenomenon of shorter duration than .1 seconds so appears to last .1 seconds. Hence the effect on the eye by a lightning flash would be about the same whether the flash lasted .1 seconds, or if it were of a thousand times greater intensity but lasting a thousandth of the time. This means that the eye would see a lightning flash about in the same manner as if its light, and so probably its energy were spread uniformly over .1 seconds.

The illumination given by a brilliant lightning flash is about of the same magnitude as good artificial illumination, perhaps one foot candle, since at night time in a well lighted room, the light of a lightning flash is still quite appreciable. Estimating roughly one watt per candle foot, a lightning flash illuminating a space of two miles square or 10^8 square feet, with one foot candle would consume 10^8 watts, and as this is the illumination as averaged by the human eye over .1 seconds, the energy is 10^7 watt-seconds, or 10,000 K. W. seconds. The energy of a large lightning flash, estimated from its light, would thus be of the magnitude of 10,000 K. W. seconds. This value, while considerable when expressed in electric quantities, is by no means so very great: reduced to heat measure, it only equals the latent heat of evaporation or condensation of about 9 lbs. of water.

As seen, an estimation of the voltage of the lightning flash from length and disruptive potential gradient of the air, does not give reasonable values, that is, the lightning flash cannot be a single discharge as that of a Leyden jar. An estimation of the voltage may then be attempted in a different manner.

Lightning flashes usually occur within thunder clouds and only rarely from cloud to cloud or from cloud to ground. They therefore seem to be rather due to equalization of potential differences within the cloud, than to discharges between oppositely charged bodies. Lightning occurs mainly when rapid condensation of moisture takes place in the air and the electric phenomena seem to be the more intense, the greater the rapidity of condensation, or rain formation. Thus the atmospheric electric disturbances seem to be connected with the condensation of water vapor to clouds and rain.

There exists normally a potential gradient in the air. That is, a potential difference exists between the air at different elevations, reaching sometimes several hundred volts per foot, so that we can estimate as a fair average, a natural potential gradient in the air, in vertical direction, of about 100 volts per foot. A point 100 feet above ground may show a potential difference of about 10,000 volts against ground. Usually the higher strata of the air are positive against the lower. The cause of this potential gradient, whether terrestrial or cosmic, is of no interest to us here, but merely its existence.

It is of interest to investigate, what effect must be expected, from our well-known physical laws, from the condensation of moisture, and agglomeration of the moisture particles to rain drops, in an atmosphere having such a potential gradient.

Assuming water vapor in a higher stratum of the atmosphere to condense to moisture particles, these moisture particles have the potential of the air in which they float, that is, have a considerable potential difference, perhaps hundred thousands of volts, against ground, and so contain an electric charge against ground. These moisture particles conglomerate with each other to larger moisture particles and ultimately

rain drops. By the collection of n^3 particles into one, the diameter of the particle has increased n fold. Its capacity has also increased n fold (the capacity of a sphere being proportional to the diameter). The particle contains, however, the accumulated charges of n^3 smaller particles, and n^3 times the charge, with n times the capacity, gives n^2 times the potential. It follows herefrom that with the conglomeration of the water particles, their potential must increase rapidly, proportionately to the square of their diameter. The conglomeration of moisture particles in the clouds is, however, very uneven, due to the uneven distribution of moisture, as is plainly seen by looking at any cloud: dense or dark parts representing considerable condensation and so considerable moisture content, alternate with light parts, in which little or no condensation occurs. As a result thereof, starting with a uniform potential in the stratum of the air, where condensation begins, differences of potential distribution by necessity result from the differences in the condensation of water vapor to moisture and the accumulation of the moisture particles to larger ones, that is, the denser portions of the cloud are at a higher potential than the lighter portions. Thus, starting with uniform potential, and thus zero potential gradient in the air at the moment of the beginning of condensation, potential differences and thus potential gradients appear.

Such potential differences in the clouds increase with increasing agglomeration of moisture particles to rain drops, and so the potential gradient rises. Assuming even as low a potential gradient as 100 volts per foot in the cloud at the beginning of agglomeration of moisture particles. the collection of n^3 such particles to one rain drop of n times the diameter and so n times the capacity, but containing the static charge of n^3 particles, gives n^2 times the potential, and since the dis-

tances between the particles are now n times as large, the potential gradient has increased n fold. That is, by conglomeration of water particles, the potential gradient rises proportionately to the diameter of the particles. Estimating then the average diameter of moisture particles as 10^{-4} inches at the beginning of agglomeration, when the potential gradient in the cloud is about 100 volts per foot, then the breakdown potential of the air, of between 100,000 and 200,000 volts per foot, would be reached when the drops have reached about .1 to .2 inches diameter, that is, the size of rain drops.

Potential gradients in the cloud thus gradually rise, until somewhere the disruptive strength of the air is reached, and a discharge passes, equalizing the voltage at this spot. This, however, causes a greater potential gradient at the ends of the discharge, exceeding the breakdown strength of the air, and so causes a second discharge, following partly over the path of the first, then a third and so on, until all of the potential differences or inequalities of the potential distribution in the cloud, are leveled down by a series of successive discharges. The phenomenon thus is similar to that of a landslide, setting off another and another landslide. Or it can best be pictured by representing the unequal moisture distribution in the cloud by a relief map built of wet sand, the dense portion of the cloud, and therefore the portions of high potential, being represented by the hills, the light or low potential portions of the cloud by the valleys of the relief map. As soon as the sand dries, somewhere, where the declivity is very steep, that is, the potential gradient is very high, a slide occurs, this causes another slide and so on, until the whole configuration of sand settles down to a flat and smooth shape, the hills are leveled off and the valleys filled.

The existence of such successive discharges, following each other after appreciable intervals of time in the same path, has been shown by the photographs of lightning flashes taken with a rotating camera. In this case, by the motion of the camera the successive flashes are recorded side by side, and sometimes more than forty successive discharges have been counted, the whole phenomenon lasting about .6 seconds, that is, quite an appreciable time.

Oscillograms of lightning discharges from (dead) transmission lines also showed the frequent occurrence of multiple strokes, or strokes following each other within a fraction of a second.

It follows herefrom, that lightning flashes in the clouds, of several miles' length, occur without any considerable potential difference between the ends of the flash, but result from the disruptive equalization of the unequal potential distribution in the clouds, caused by unequal vapor density and so unequal condensation and conglomeration of moisture particles.

This also explains the relatively small tendency to discharges between cloud and ground, across a space in which no condensation takes place and so no unequal potential distribution supplies the power of the discharge: although the distance between cloud and ground is smaller than the distance traversed by a lightning flash in the clouds, and the average potential difference between cloud and ground probably is greater than the potential differences in the clouds, a discharge to ground probably occurs in general only where by a heavy downpour of rain a range of high potential is carried bodily part ways down to ground. This also may explain, that lightning discharges to the ground are usually followed by a heavy downpour of rain.

The potential gradient in the air may rise to disruptive values in still another, slightly different manner, and lead to lightning discharges without being accompanied or followed by rain. By conglomeration of moisture particles the potential gradient rises, as described above, but before the water drops have reached sufficient size to precipitate as rain, evaporation again sets in: for instance by the drops falling to a lower and warmer stratum of the air, or by intercepting the heat of the sun's rays, and the drops thus dwindle away. The decrease in size of the drops represents a decrease of capacity, the capacity being proportional to the diameter, and as each drop retains the same charge, its potential increases with the decrease of size, without limit, and so also the potential gradient until its disruptive value is reached and the lightning discharge occurs. This phenomenon is frequently observed towards the evening of a hot summer day, and is called "heat lightning", and, being the result of evaporation, thus does not lead to rain.

Estimating then as disruptive strength of air under discharge conditions in a non-uniform field, and at the reduced air pressure in the clouds, 100,000 volts per foot, the average potential gradient in the path of the lightning discharge through the clouds would be about 50,000 volts per foot. This gradient, however, would not be unidirectional, but the potential would rise from a low, or even negative value at a light portion of the cloud, to a maximum value at a dense position, then decrease again, that is, give a gradient in opposite direction, to a light position, etc., and the potential gradient would vary from nothing at a maximum potential point, to a maximum, equal to the breakdown strength of air at the starting point of the discharge, to zero at a minimum potential point, etc.

To estimate the current which discharges in the lightning flash, the conductivity of air in the path of the discharge, and the diameter of the discharge are required, and as both are unknown, any estimate must be very approximate only. The specific resistance of gases and vapors decreases with increasing temperature and with decreasing pressure. It is a few ohm centimeters at atmospheric pressure and the high temperature of the magnetite or carbon arc, and is also a few ohm centimeters at the low temperature and low pressure of a high current Geissler tube discharge. The mercury arc stream also gives a specific resistance of a few ohms. The temperature of the air in the lightning discharge probably is moderately high, but the pressure is also not far from atmospheric, so that 100 ohm centimeters may not be very far from the true magnitude of the resistance. Estimating one to two feet as the diameter of the discharge path, and 100 ohm centimeters as the specific resistance, and allowing for the inductance, gives, with an average potential gradient of 50,000 volts per foot, a current of about 10,000 amperes.

The heating effect and the magnetic effect of lightning strokes also point to the passage of currents of some thousand amperes.

Assuming then the average potential gradient in the lightning flash as 50,000 volts per foot, the current as 10,000 amperes, a lightning flash of two miles' length would represent a power of 5×10^9 K. W. Estimating the energy of the discharge, as approximated from the photometric consideration, as 10,000 K. W. seconds, the duration of the discharge would be: $10^4 / 5 \times 10^9 = 2 \times 10^{-6}$ sec., or two-millionths of a second.

The discharge probably is oscillatory. In view of the high resistance of the discharge path, the damping effect must be very great, that is, a very large part or nearly all the energy

is expended in the first half-wave; that is, the discharge consists of only one or very few half-waves. With a duration of the discharge of 2×10^{-6} seconds, assuming two half-waves as an average, gives 500,000 cycles.

The frequency of oscillation of the lightning flash thus appears to be of the magnitude of half a million cycles.

Since the velocity of propagation of electric disturbances is the velocity of light, or 188,000 miles per second, the wave length of a discharge of 500,000 cycles is $\frac{188,000}{500,000} = \frac{3}{8}$ miles, or about 2000 feet.

A wave length of 2000 feet means that the current in the discharge flows in one direction for 1000 feet, in the opposite direction, that is, with opposite potential gradient, in the next thousand feet, etc. That is, in our former discussion, the average distance through which the potential gradient has the same direction, or the distance between maximum and minimum, between densest and lightest parts of the cloud is about 1000 feet. This agrees fairly well with the appearance of the clouds to the eye, and it also agrees in magnitude with the distance over which the wind velocity varies, in gusts, as shown by Prof. Langley in his investigation on the "internal energy of the wind".

It appears herefrom, that the varying wind velocity as measured by Prof. Langley, that is, the gusty character of the air currents, results not only in an internal mechanical energy, which the bird utilizes for soaring, but also results in unequal moisture distribution, and so, when condensation occurs, in an "internal electrostatic energy" of the thunder cloud, which discharges as lightning.

With an average length of the half-wave of 1000 feet, and 50,000 volts per foot as potential gradient, the potential

differences in the clouds would be of the magnitude of fifty million volts. These are values which appear reasonable.

Assuming that a lightning flash drains the electric energy of the cloud within a radius of about 100 to 200 feet from the path of the discharge, this affords a different method of estimating the magnitude of the energy of the lightning flash: assuming for instance saturated air at 40°C mixing with air at 0°C, condensation of a part of the moisture occurs which can easily be calculated. Assuming that this moisture has conglomerated to rain drops of .1" to .2" diameter, the number of such drops in a space of two miles' length, and 200 to 400 feet diameter, can be calculated, and also their electrostatic capacity. With a wave length of 2000 feet, and a potential gradient of 50,000 volts per foot, from the capacity follows the energy of the electrostatic charge, which discharges as lightning flash. This is found under the above assumption, as of the magnitude of 10,000 K. W. seconds, so agrees with the results derived from the photometric considerations.

To conclude then, as approximate values of magnitude of the electric quantities in a lightning flash may be estimated:

Average potential gradient: 50,000 volts per foot at the moment of discharge.

Average potential difference between different points of the cloud: 50 million volts.

Average current in the discharge 10,000 amperes.

Average duration of the discharge $\frac{1}{500,000}$ sec.

Average frequency of discharge: 500,000 cycles.

Average energy of the discharge: 10,000 K. W. sec., or seven million foot pounds.

II. LIGHTNING IN ELECTRIC CIRCUITS.

Of greatest importance to an electrical engineer are the high potential phenomena produced in electric circuits by atmospheric lightning as well as by other causes, frequently internal to the circuit, which give the same or similar effects to such an extent, that it has become customary when dealing with electric circuits, to distinguish between external or atmospheric lightning, and internal lightning, as caused by electric circuit disturbances or defects, such as sudden changes of load, or arcing grounds, etc.

While a very large amount of data on high potential phenomena in electric circuits has accumulated, the possible variety of phenomena is so great that an intelligent understanding of the phenomena, as is required for effective protection of the circuits, is feasible only by a theoretical investigation of the high potential phenomena which may be expected in electric circuits, and a comparison thereof with the observed effects.

In general, the high potential phenomena possible in electric circuits are the same three classes of phenomena which can occur in any medium, as a body of water, which is the seat of energy.

1. Steady electrostatic stress, that is, a gradual rise of potential of the total circuit against ground, until a discharge occurs somewhere; just as in a body of water, as a river, the pressure, that is, the water level, may gradually rise, until it breaks through the embankment.

2. Impulses, or traveling waves, similar to the ocean waves rolling over the surface of the water.

3. Standing waves, or oscillations or surges, similar to the oscillation of a tuning fork, or a violin string.

A more extended discussion on the three forms of electric disturbances, and their causes, is given in a paper read before the A. I. E. E.*

Steady electrostatic stress obviously can occur only where the circuit is very well insulated from the ground, but not in a grounded circuit, or a leaky circuit, as low voltage circuits usually are, and such static stresses can be eliminated by a permanent leak, that is, a high resistance connection between the circuit and the ground.

As sources of impulses or traveling waves only two characteristic phenomena may be considered here: the lightning flash, or induction by the clouds, as external, and the arcing ground as internal cause.

Assuming a thunder cloud to pass over the line. The ground below the cloud then assumes an electrostatic charge, corresponding to the opposite charge of the cloud. The transmission line, as part of the ground, thus also assumes a static charge, higher than that of the ground, since it projects above it. Any equalization of the potential distribution in the cloud by a lightning flash, as discussed in the preceding, requires a change in the electrostatic charge of the line, corresponding to the changed potential difference between ground and cloud above the ground, and the static charge thus set free on the line rushes as an impulse or wave along the line. The wave shape of such impulses induced by cloud discharges is in general not a smooth sine wave, but may be very irregular: during the equalization of the cloud potential by the lightning flash, the potential difference against ground, of the part of the cloud above the electric circuit, may vary in almost any conceivable manner, thus giving rise to very different wave shapes of the impulses. So some impulses may rise very rapidly, with

*A. I. E. E. Transact. March, 1907: "Lightning Phenomena in Electric Circuits."

extremely steep wave front, and slowly die down. Others may rise slowly, then suddenly fall and reverse, or a series of oscillations may occur in the impulse, etc. If the lightning flash is parallel with the line, simultaneous impulses of different directions may be produced, corresponding to the different directions of the potential gradient in the different parts of the lightning flash, and these waves, of different directions, intensity and wave length, traveling over each other, then produce a very complex system of phenomena. So for instance, by the interference of two impulses of nearly equal wave length, moving in opposite directions, a high voltage point may be produced, traveling slowly along the line, and visible to the eye as a luminous streak.

The frequencies of these impulses then are those corresponding to the frequencies of cloud discharge, that is, of the magnitude of hundred thousands of cycles per second. With the velocity of light, 188,000 miles per second, they travel along the line, until they gradually fade out by the dissipation of their energy, or are reflected at an open end of the line, or at the entrance to the station are broken up by partial reflection, in reactances, and interference between the reflected waves, the incoming waves and the waves passing over the reactances, and so give rise to systems of standing waves or oscillations, similarly as an ocean wave rolling on to a sloping beach breaks up into surf.

Where a traveling wave is reflected, the combination of the reflected wave and the incoming wave produces a standing wave or oscillation, that is, a wave in which the voltage maxima and the zero points or nodes have fixed positions on the line.

By superposition of the wave maxima of incoming and reflected wave, the standing wave rises to a maximum double

that of the traveling wave. Where different oscillations or standing waves superimpose upon each other, their maxima subtract at some places and add at others, and thus again double the voltage, that is, a traveling wave or impulse, breaking up into systems of oscillations at a station, doubles and quadruples the potential; so that a traveling wave of moderate potential may cause dangerous voltages when breaking up into oscillations, just as in the ocean surf, the waves rise to far greater heights than in the on-rolling ocean wave before it reaches the beach.

If we consider that the impulses traveling along the line are not sine waves, but of very irregular shape, that is, can be considered as consisting of a fundamental of some hundred thousand cycles, and numerous higher harmonics of still greater frequency, and each of the components when breaking up at the station gives rise to a set of oscillations at every interference point, that is, at every reactance, the complexity of the phenomenon can be imagined.

Since the equalization of cloud potential usually occurs by a series of successive discharges in short intervals, a small fraction of a second, and each discharge gives rise to an impulse in the line, and so a system of oscillations at the station, whatever protective device is used, must restore itself instantly after a discharge, so as to receive the next following discharge. Any device depending on mechanical motion to restore itself after a discharge to operative position, therefore fails to protect, when a series of discharges follow each other in very rapid succession, as discussed above.

Traveling waves very similar in character to those due to induction from the clouds, but frequently of far greater volume, sometimes occur in an electric circuit from internal causes, as arcing grounds, or spark discharges.

Let, for instance, a spark occur in an insulated underground cable system between one of the conductors and the grounded cable armor, through a weak spot in the insulation, as a faulty joint or a cable bell. Normally a potential difference exists between the cable conductor and the ground, equal to the Y potential of the system, and so an electrostatic charge on the conductor corresponding thereto. A spark passing between conductor and ground, connects it to ground, and the charge of the conductor so passes over the spark as arc to ground. As soon, however, as the conductor is discharged and at ground potential, the arc between conductor and ground ceases, since there is no voltage left to maintain it, and so the conductor disconnects from ground. The conductor then charges itself again to its normal Y potential and during the in-rush of the charge, momentarily the potential builds up to double voltage. Thereby a spark again passes between conductor and ground, discharges it again, opens after discharge, again causes a spark to pass, etc. So a series of successive sparks occur between conductor and ground, discharging the conductor by currents which momentarily rise to very high values, the discharge current of the capacity of the conductor against ground, over a path of practically no resistance. Each spark discharge sends out an impulse or traveling wave, and thus a spark discharge between conductor and cable armor, or in the same manner an arcing ground on an overhead transmission line, as is for instance caused by a broken insulator, produces a continuous series of impulses or traveling waves, which follow each other with the rapidity of charge and discharge of the cable or the line, that is, many thousands per second, and so give what has been called a *recurrent surge*. In a long distance transmission line, the frequency of the recurrent surge usually

is somewhat lower than in an underground cable system, but is still thousands of impulses per second.

The frequency of oscillations occurring in electric circuits varies over an enormous range: from low frequencies, very little above alternator frequency, up to hundreds of millions of cycles per second; and the effect of the oscillations in the system therefore varies accordingly: from the relatively harmless static displays; brush discharges, streamers, sparks, etc., of extremely high frequencies, down to the disastrous high power low frequency short circuit oscillations, in which even in 10,000 volt systems, currents of many thousands of amperes may surge, which voltages approaching 100,000, and with which no protective device can cope, which does not have unlimited discharge capacity, that is, contains no resistance whatever in the discharge path.

III. LIGHTNING PROTECTION OF ELECTRIC CIRCUITS.

From the preceding considerations it follows that the problem of protecting electric circuits from lightning is twofold:

1. To guard against high potential disturbances entering the circuit from the outside or originating in the circuit.
2. To discharge harmlessly to ground, whatever high potential phenomena may appear in the circuit.

From atmospheric electric disturbances, complete protection can be secured by putting the circuit under ground, or, where this is not feasible, to put the ground over the electric circuit. This means the use of grounded overhead wires. The overhead ground wires so protect the circuit the more completely, the more they realize a complete shield interposed between line and sky. While complete protection thus would

require a system or network of grounded conductors above, beside, and also below the transmission line, very good protection in most cases is secured by a single ground wire of good conductivity, installed well above the line; and in no place of electric transmission systems can money be more efficiently spent, than in securing good overhead ground wire protection.

To guard against the appearance of internal lightning requires constant watchfulness in the design, construction and operation of the system, to avoid all conditions which may lead to the formation of oscillating arcs. Thus poor contacts, loose joints, masses of insulated metal near high potential conductors, etc., should be carefully avoided.

The disturbances which have to be taken care of by the lightning arresters proper, are steady accumulation of static pressure; impulses or traveling waves; oscillations or surges; occurring singly or in groups, and of frequencies varying between many millions of cycles and ordinary machine frequencies; and recurrent surges, that is, impulses and oscillations, usually of high frequency, following each other in very rapid succession, usually thousands per second.

It is necessary that the discharge over the lightning arrester should occur with the least possible disturbance to the system, that is, the discharge current should be as small as permissible without causing a voltage rise due to the resistance of the discharge path. At the same time, the protective devices must be able to discharge practically unlimited currents, that is, currents of the magnitude of the momentary short circuit current of the system. This obviously requires that the protective devices should have no appreciable resistance in the discharge path. Any lightning arrester containing series resistance obviously fails to protect as soon as the discharge current is so large that the ohmic drop across the resistance

becomes serious, and the maximum discharge current which may occur, is the short circuit current of the system, that is, extremely large.

Three types of protective devices are at present available.

1. The circuit is connected to ground by a single spark gap set for a voltage exceeding the normal operating voltage by a safe margin: the so-called horn gap, or goat horn lightning arrester. As soon as the voltage rises beyond the value for which the spark gap is set, it discharges, and the system is short circuited to ground, until the arc rises and gradually blows itself out. As this requires an appreciable time, motors and converters have usually dropped out of step, and the generators have broken synchronism, that is, the system is shut down and has to be started up again. This type of protection therefore is not particularly favored in systems which require reasonable continuity of service, but if used, it is considered rather as an emergency device in addition to other arresters and is then adjusted for much higher discharge voltage. A reduction of the current over the horn gap by series resistance is not permissible, since it correspondingly reduces the protective value, as explained above, and the arrester ceases to protect against a high power surge. While such surges are relatively infrequent, their destructiveness is such that protection against them is especially needed. Fuses in series with the horn gap, if they open slowly, would still shut down the system, and if opening very rapidly, the shock of the explosive opening of the fuse on the short circuit current of the system may be disastrous. Obviously, the use of series fuses require a multiplicity of spark gaps to give continuity of protection.

2. The type of lightning arrester now almost universally used is the multi-gap arrester, which short circuits the system for one-half wave only. It consists of a large number

of spark gaps between metal cylinders, in series with each other. As now designed, different sections of the gaps are shunted with different resistances, for the purpose of affording equal protection against all frequencies, and adjusting automatically the resistance of the discharge path to the volume of the discharge, as for instance, discharge slow accumulations of potential over a very high resistance, short circuit surges over a path of zero resistance, and thus pass a discharge with the minimum shock on the system. The operation of the multi-gap—which by the way is suitable only for alternating current systems—depends on the non-arcing character of certain metals. Metals of low boiling point, as mercury or zinc, cannot maintain an alternating current arc, but the arc goes out when at the end of the half wave, the current falls to zero, and a very much higher voltage is required to again start an arc for the next half-wave.* Alloys of such metals, usually zinc, with metals of high melting point, as copper, are therefore used as terminals in the multi-gap arrester.

A discharge over the multi-gap arrester short circuits the system for the rest of the half-wave during which the discharge occurs. At the end of the half-wave, the current falls to zero, and the reverse current cannot start, that is, the circuit of the arrester is opened.

A short circuit on the system, for a fraction of a half-wave, does not interfere with the operation of synchronous apparatus, that is, the operation of the system is not affected by a discharge over the multi-gap arrester.

In a large system, the short circuit current is very considerable; its power, and thus the heating effect produced by it, is enormous. The energy, and thus the heat produced by the short

See paper A. I. E. E. Transact. 1906, p. 789. "Transformation of Electric Power into Light."

circuit current during the fraction of the half-wave, which the discharge over the multi-gap arrester lasts, is moderate, due to its very short duration, and can easily be absorbed and radiated by the arrester; so that even if lightning discharges rapidly follow each other for some time, they can be taken care of by the arrester with moderate temperature rise: assuming a vicious thunder storm, in which lightning flashes succeed each other practically continuously, several per second. Each discharge causes a short circuit over the lightning arrester, varying in duration from nearly a half-wave—if the discharge occurs at the beginning of a half-wave—to practically nothing—if the discharge takes place near the end of a half-wave—that is, in average, for one-half of one-half wave, or $\frac{1}{2 \times 40}$ sec., in a 60 cycle system. Therefore from two to three lightning discharges per second would still short circuit the system over the multi-gap arrester only for 1% of the total time, and the heating effect, caused by a short circuit during 1% of the time, can be taken care of by the arrester for a considerable period.

Let us see, however, what happens to the multi-gap lightning arrester in case of the appearance of a recurrent surge, as an arcing ground, that is, discharges following each other in rapid succession, thousands per second. The first discharge, passing over the lightning arrester, short circuits the system for the rest of the half-wave, and at the end of the half-wave, the arrester functionates properly, that is, opens the circuit. At the next moment, however, at the beginning of the next half-wave, the next oscillation of the recurrent surge again discharges over the arrester, and thus again short circuits. That is, with a recurrent surge, the multi-gap arrester at the end of every half-wave opens the circuit, at the beginning of the next half-wave, the next oscillation of the recurrent surge short circuits again. As far as the effect on the operation of

the system, and the heating of the arrester is concerned, a recurrent surge causes a permanent short circuit on the system, except that at the beginning of every half-wave, for a short period, the circuit is opened and free for the appearance of disruptive voltages elsewhere, and so apparently, simultaneous with the short circuit, destructive high potentials may appear in the system. The heating effect of the short circuit current, which occurs at every half-wave, rapidly destroys the arrester. In such cases, to save the arrester, it has been customary to insert a series of auxiliary gaps, which are thrown in by the blowing of a fuse shunting them, and raise the discharge voltage of the arrester so that the recurrent surge does not pass over it. It is obvious, that in this case the arrester ceases to protect the system against the recurrent surge: but if left in circuit, the destruction of the arrester would put it out of operation anyway.

It is obvious now, that no lightning arrester, which functionates by short circuiting the system for the rest of the half-wave, during which a discharge occurs, can take care of and protect against a recurrent surge, since the proper functioning of the arrester, with a recurrent surge, represents a permanent short circuit on the system over the arrester, and so a destruction of the arrester, no matter whose make it may have been, and a shutdown of the system.

3. To take care of a recurrent surge, a protective device would thus be required, which does not short circuit the system even for one half-wave, but which never allows the normal voltage of the system to pass a current over the arrester, but acts as a short circuit for any excess voltage above the normal voltage. The possibility of such a device we can understand by considering the effect, which in a direct current circuit a storage battery would have, when shunted between the

circuit and the ground. Assuming for instance in a 500 volt trolley circuit, a 500 volt storage battery of very high capacity, that is, negligible internal resistance, permanently connected between line and ground. With the normal line potential of 500 volts, no current would pass over the battery to ground, except the very small current required to maintain the battery charged. No rise of voltage, however, could occur in the system by lightning or any other cause, since any voltage above 500 volts, the counter e. m. f. of the battery, would be short circuited to ground through the battery, and such a battery would thus give perfect protection against any high voltage disturbances in the system. In case of a recurrent surge, the current discharging over the battery would be the short circuit current of the excess voltage, that is, the surge potential, and the heating effect of this current is negligible, since high potential high frequency phenomena are of limited power and especially of limited current, as condenser discharges.

A storage battery obviously is not suitable for alternating current and would not be practical in any case, as it requires a cell for every two volts. The same effect, however, is produced at a much higher voltage, in an alternating current circuit, by the aluminum cell. If such a cell, consisting of two aluminum plates in certain electrolytes, is exposed to an alternating voltage, a film forms on the aluminum plates, which holds back the impressed voltage, that is, acts like a counter e. m. f. equal to the impressed e. m. f., so that practically no current passes through the cell, or only the small current required to maintain the film, of a magnitude of about .01 amperes per square inch plate surface, while for any sudden rise of voltage the cell acts as a short circuit for the excess voltage. Over the storage battery, the aluminum cell has the advantage of higher voltage: a single cell can take care of

300 to 400 volts and even more, and also that it does not have a fixed counter e. m. f., but a counter e. m. f., which adjusts itself to equality with the impressed voltage, at any value up to about 600 volts per cell. Assuming for instance an aluminum cell connected across an alternating e. m. f. of 300 volts. With the film formed, a negligible current passes through the cell, for instance, $\frac{1}{4}$ of an ampere, maintaining the integrity of the film. If now the voltage is suddenly raised to 330 volts, in the first moment the cell acts as a short circuit of the excess voltage, in this case, 30 volts, and for an instant a very large current, possibly hundreds of amperes if the supply source is capable of giving such a current, rushes through the cell. This current very rapidly decreases, by the film of the aluminum plates forming for higher voltage, so that in a few seconds the current is already small, and in a few minutes the normal current of 1-4 ampere again passes, but now at 330 volts impressed, and the film has formed to a counter e. m. f. equal to this higher voltage, probably has thickened. If now we again lower the voltage suddenly to 300, in the first moment the current in the cell practically disappears, and then gradually rises again, and after a few minutes is again normal at $\frac{1}{4}$ ampere, that is, the film has built down again to 300 volts. In this manner the aluminum cell adjusts its counter e. m. f. to changes of impressed voltage, by the film building up or building down. This adjustment, for moderate voltage variation, as may be expected when varying the generator voltage of the system, is quite rapid, most of the change occurring within less than a second, but is still extremely slow compared with the rapidity of lightning phenomena, and for lightning phenomena the aluminum cell therefore acts as a short circuit of the excess voltage above the normal machine voltage. Thus the recurrent surge, with a system of aluminum cells in series with

each other connected directly across the circuit, cannot produce any rise of voltage, but the excess voltage over the normal, or the surge potential, is short circuited through the aluminum cell, so causing a small increase of the current in the cells, by the superposition of the high frequency surge current over the normal leakage current of the cell, but no rise of voltage. Since the recurrent oscillations are intermittent, obviously the film of the aluminum cells cannot build up to their voltage, but remains corresponding to the machine voltage, that is, the aluminum cell can permanently discharge a recurrent surge without any short circuit of the main voltage, or any disturbance on the system.



3 8482 00426 3493

SEP 11 1987

621.3 S82g c.2
Steinmetz, Charles Proteus,
General lectures on
electrical engineering /

University Libraries
Carnegie-Mellon University
Pittsburgh, Pennsylvania 15213

Desacidified using the Bookkeeper process.
Neutralizing Agent: Magnesium Oxide
Treatment Date:

ÆTHERFORCE

UNIVERSAL
LIBRARY



138 402

UNIVERSAL
LIBRARY

ÆTHERFORCE